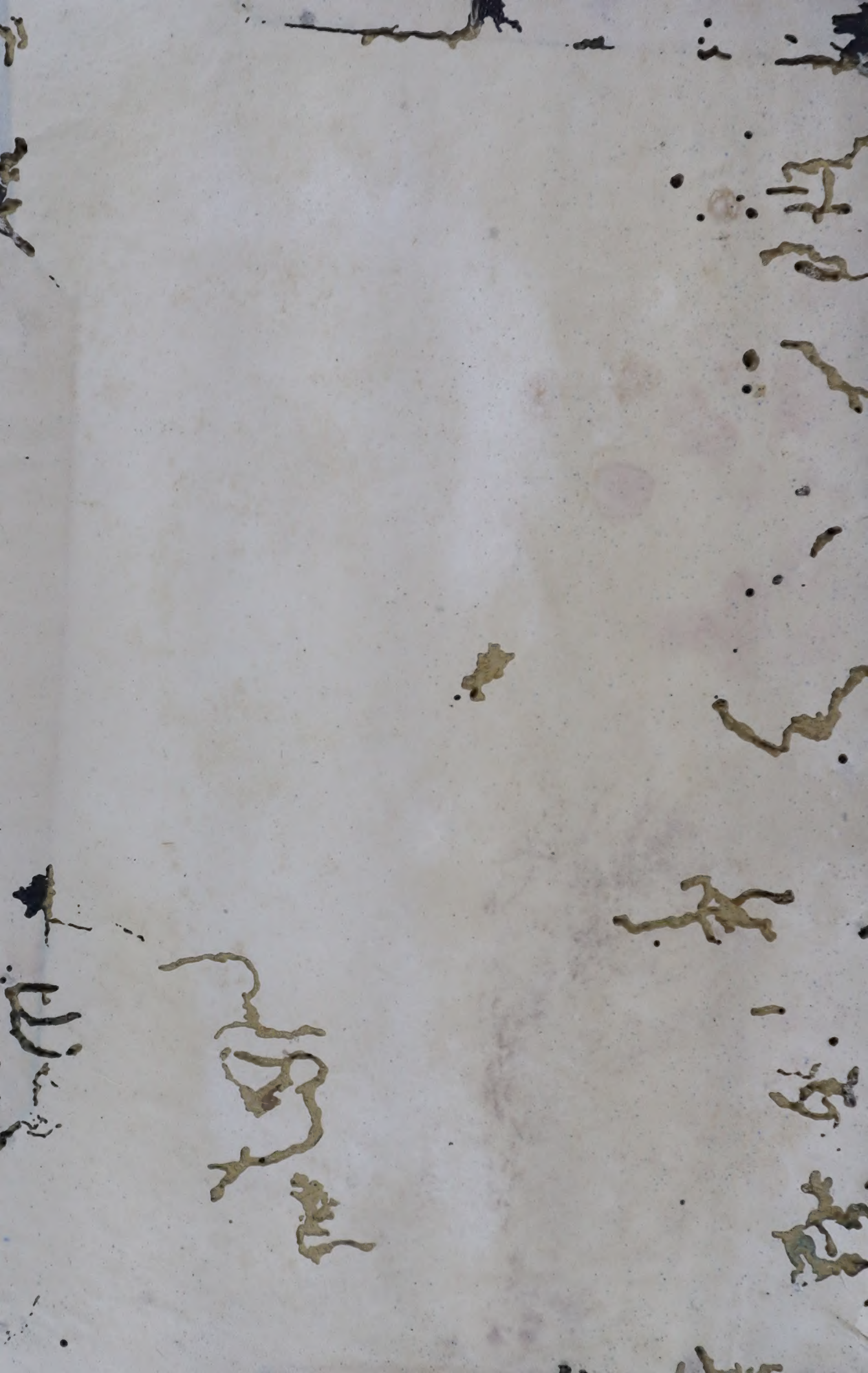




PROCEEDINGS
OF THE
ROYAL SOCIETY OF LONDON

SERIES A. MATHEMATICAL AND PHYSICAL SCIENCES

VOL 202



LONDON

Published for the Royal Society by the
Cambridge University Press
Bentley House, N.W.1

22 August 1950

*Printed in Great Britain at the University Press, Cambridge
(Brooke Crutchley, University Printer)*

*and published by the Cambridge University Press
Cambridge, and Bentley House, London*

Agents for Canada, India, and Pakistan: Macmillan



CONTENTS

SERIES A VOLUME 202

No. A 1068. 22 June 1950

	PAGE
The work of the Discovery Committee. By N. A. Mackintosh. (Plates 1 to 4)	1
Studies on the ionization produced by metallic salts in flames. II. Reactions governed by ionic equilibria in coal-gas/air flames containing alkali metal salts. By H. Belcher and T. M. Sugden	17
Transient source, doublet and vortex solutions of the linearized equations of supersonic flow. By W. J. Strang	40
Transient lift of three-dimensional purely supersonic wings. By W. J. Strang	54
The instability of liquid surfaces when accelerated in a direction perpendicular to their planes. II. By D. J. Lewis. (Plates 5 to 11)	81
A note on three-dimensional turbulence and evaporation in the lower atmosphere. By D. R. Davies	96
Resistivity of pure metals at low temperatures. I. The alkali metals. By D. K. C. Macdonald and K. Mendelssohn	103
The fine structure of the hydrogen α line. By H. Kuhn and G. W. Series. (Plate 12)	127
The aerodynamic drag of grassland. By F. Pasquill	143

No. A 1069. 7 July 1950

The molecular orbital theory of chemical valency. III. Properties of molecular orbitals. By G. G. Hall and Sir John Lennard-Jones, F.R.S.	155
The molecular orbital theory of chemical valency. IV. The significance of equivalent orbitals. By Sir John Lennard-Jones, F.R.S. and J. A. Pople	166
The kinetics of the interaction of atomic hydrogen with olefines. V. Results obtained for a further series of compounds. By H. W. Melville, F.R.S. and J. C. Robb	181
Statistical mechanics of the two-dimensional assembly. By H. N. V. Temperley	202
Travelling disturbances in the ionosphere. By G. H. Munro	208
Adhesion of solids and the effect of surface films. By J. S. McFarlane and D. Tabor	224
Relation between friction and adhesion. By J. S. McFarlane and D. Tabor	244
Homotopy invariants and continuous mappings. By Su-Cheng Chang	253
The $\text{CH}_2 : \text{CH} \cdot \text{CH}_2 - \text{CH}_3$ bond dissociation energy and the heat of formation of the allyl radical. By A. H. Sehon and M. Szwarc	263
Diffraction by a plane screen. By E. T. Copson	277
The structure of linear relativistic wave equations. I. By K. J. Le Couteur	284

No. A 1070. 7 August 1950

	PAGE
On the stochastic theory of continuous parametric systems and its application to electron cascades. By H. J. Bhabha, F.R.S.	301
The molecular orbital theory of chemical valency. V. The structure of water and similar molecules. By J. A. Pople	323
The molecular orbital theory of chemical valency. VI. Properties of equivalent orbitals. By G. G. Hall	336
Finite strains in an anisotropic elastic continuum. By J. G. Oldroyd	345
Contributions to the theory of heat transfer through a laminar boundary layer. By M. J. Lighthill	359
The conductivity of thin wires in a magnetic field. By R. G. Chambers	378
The structure of linear relativistic wave equations. II. Representations. By K. J. Le Couteur	394
Finite elastic deformation of incompressible isotropic bodies. By A. E. Green and R. T. Shield	407
The quaternary system aluminium-iron-cobalt-nickel, with reference to the role of transitional metals in alloys. By G. V. Raynor and M. B. Waldron	420
Identification of microseismic activity with sea waves. By J. Darbyshire	439

No. A 1071. 22 August 1950

Constitution and mechanism of the selenium rectifier photocell. By J. S. Preston	449
The Bloch integral equation and electrical conductivity. By P. Rhodes	466
Degree of polymerization and chain transfer in methyl methacrylate. By Sadhan Basu, Jyotirindra Nath Sen and Santi R. Palit	485
Electronic levels in simple conjugated systems. I. Configuration interaction in cyclobutadiene. By D. P. Craig	498
Asymptotic expansions for the hypergeometric functions occurring in gas-flow theory. By T. M. Cherry	507
Resistivity of pure metals at low temperatures. II. The alkaline earth metals. By D. K. C. MacDonald and K. Mendelssohn	523
Term values in hybrid states. By W. Moffitt	534
Aspects of hybridization. By W. Moffitt	548
Surf beats: sea waves of 1 to 5 min. period. By M. J. Tucker	565
Programme organization and initial orders for the EDSAC. By D. J. Wheeler	573
Index	590

The work of the Discovery Committee

By N. A. MACKINTOSH, C.B.E., D.Sc., *Director of Research, Discovery Investigations*

(Lecture delivered 26 January 1950—Received 16 February 1950)

[Plates 1 to 4]

The geographical field in which most of the Discovery Committee's work has been carried out during the past 25 years is the Southern Ocean. This zone of continuous deep water, very rich in marine life, supports one major industry—the whaling industry—but is otherwise little developed as yet, and seldom visited. It is not easy to find a short descriptive label for the work itself, but nearly all of it comes under the headings of deep-sea oceanography, whales and whaling, or Antarctic geography, and much of it is concerned with the interrelations of these subjects.

Since the beginning in 1924 the Discovery Committee has worked under the Colonial Office, but in 1949 the Committee's functions, together with the scientific staff, the ships, and other assets, were taken over by the Admiralty, and now form part of the new National Institute of Oceanography. The Discovery Committee, in its original form, has been dissolved, but it is encouraging to know that the continuation of its work is assured.

ORIGIN AND PROGRAMME OF THE DISCOVERY COMMITTEE

The Committee was appointed in 1924 by the Secretary of State for the Colonies, to carry out the recommendations made in a *Report of the Interdepartmental Committee on Research and Development in the Dependencies of the Falkland Islands* (Cmd. 657, 1920). The waters of the Dependencies form a substantial sector of the Southern Ocean, and, at that time, constituted the world's most important whaling centre. The *Report* reviewed the natural resources of the region, and in view of the rapid development of the whaling industry, it was proposed that the investigations should bear largely on the biology of whales and on their oceanic environment, mainly with a view to the proper regulation of the industry; but it urged the importance of research on very broad lines and recommended the study of the hydrography and natural history of the whole area.

The Discovery Committee was an Executive Committee, acting on behalf of the Government of the Dependencies but responsible to the Secretary of State. The first chairman was Mr E. R. Darnley of the Colonial Office, and among members were representatives of the Admiralty, the Ministry of Agriculture and Fisheries, the British Museum (Natural History), and the Royal Geographical Society. Dr Stanley Kemp was appointed Director of Research, and the work was financed from a substantial 'Research and Development Fund', built up by the Government of the Falkland Islands from taxes on whale oil processed in the Dependencies. Captain Scott's former ship, the *Discovery* (figure 8, plate 3), was purchased and refitted for oceanographical research, a Marine Station was built at South Georgia

in the Dependencies (figure 4, plate 1), and in 1925 work in the field began as an expedition called the 'Discovery Expedition'.

The programme was to include research on the general biology of whales by examination of the carcasses brought into whaling stations, and on their distribution and movements by the attachment of marks to the living whale at sea. At the same time the factors governing the distribution of whales, and, indeed, the whole marine fauna and flora, were to be studied by means of a comprehensive oceanographical survey. Broadly these lines of approach are in accordance with the methods of modern fisheries research, but the expedition had to adapt them to a new object in a new field, and to develop certain new technical methods appropriate to long voyages in severe weather conditions, to intensive deep-sea observations, and to such large and inaccessible animals as whales. It had further been decided that the oceanographical work should not be confined to subjects of immediate economic application, and that general deep-sea research, in the tradition of the Challenger Expedition, should be included. This was an ambitious programme, and it needed a man of the calibre of the late Dr Kemp to put it into operation. The Discovery Committee was a strong committee which formulated a policy and guided the work, and in the first place, I believe, it owed its existence to the initiative of Mr Darnley; but the work of building up the Discovery Investigations devolved primarily upon Dr Kemp. He relinquished his post in 1936 to become Director of the Plymouth Laboratory but continued as a member of the Discovery Committee until his death in 1945, and at every step we are still conscious of the effects of his wise judgement and foresight.

The Committee's work in the first place was to launch an expedition, and thereafter there was little to be done in England until the first voyages of the ships were completed. It may therefore be best to give an account first of the work in the field, the methods employed, and the principal results, and to deal later with the organization which was built up in London and the arrangements for research at home.

THE 'DISCOVERY EXPEDITION', 1925-1927

The Dependencies of the Falkland Islands lie between the meridians of 20° and 80° W. (figure 1), and they include part of the Antarctic continent and several scattered islands and groups of islands separated by a few hundred miles of deep ocean. Most of them are outside the Antarctic circle, but they are all mountainous and heavily glaciated. Whaling was carried out from land stations at South Georgia and moored factory ships in the South Shetland Islands and South Orkney Islands, but it has now ceased at these shore bases except at three of the stations at South Georgia.

The Marine Station with its laboratory was built close to one of the whaling stations at South Georgia (plate 2), and work began here early in 1925. Each whale brought into the station was measured and examined, and we were interested mainly in their breeding, growth, and age, their food, and the range of variation in their dimensions and external characters. The observations made were partly new, and partly based on methods initiated by other biologists, but the principal object was

to get data from a sufficient number of whales for quantitative treatment. Good progress was made, and by April 1927 we had examined nearly 1700 whales. With some changes of staff the station was kept open until 1931, by which time we had



FIGURE 1

the 'dossiers' of about 3700 whales. I shall return to the results of all this work later. In the meantime the Royal Research Ship *Discovery* undertook a preliminary survey of the South Georgia whaling 'grounds' in 1926 and was later joined by the R.R.S. *William Scoresby*. More thorough ocean surveys of the whaling grounds were then carried out around South Georgia and the South Shetland Islands, and

lines of deep-sea 'stations' were extended over the Falklands sector of the Southern Ocean and parts of the South Atlantic. In 1927 the *Discovery* returned to England, but in the next two years much of this work was repeated and extended by the *William Scoresby*. It can be looked upon now mainly as a preliminary, but intensive, survey of one part of the Southern Ocean, and it revealed the general content and distribution of the plankton population, and the local distribution of water masses and currents.

THE *DISCOVERY II* AND EXPANSION OF THE FIELD

The Committee had not been in a position at first to plan a scheme of research extending over any long period, and in 1927 the future of the work was still very uncertain. In the following year, however, the *Discovery* was taken over for the British, Australian and New Zealand Antarctic Expedition, and it was then decided that a steel steamship, the R.R.S. *Discovery II*, should be built, with greater range and power, and specially designed for deep-sea research.

It was the acquisition of this ship in 1929 which converted an expedition into a long-term service. It changed the whole complexion of the work, and the title 'Discovery Investigations' became more in keeping with its nature. For a proper understanding of the Discovery Committee's work I think it is essential to regard it as something intermediate between an expedition and a permanent research institution. In any new venture results are achieved, at least in part, through trial and error, and although an expedition may make new discoveries of great importance, it is often a matter for regret that when it is over, hard-won experience is lost through the dispersal of the staff to other occupations, and valuable equipment, often capable of long-continued use, is disposed of. The Discovery Investigations have evolved in some ways perhaps in a haphazard, hand-to-mouth manner, with some of the adventures, mistakes and new achievements of an expedition; but the strength of the present organization lies in the fact that after breaking new ground it has been able to accumulate experience and retain its capital equipment, and that it can count on the help and co-operation of many scientists who have been associated with it.

After the *Discovery II* was commissioned the field of work expanded. The preliminary work of the old *Discovery* and the *William Scoresby* had shown that although local oceanographical conditions could be measured and mapped, they could not be properly understood without fuller knowledge of the water masses and currents of the whole circle of the Southern Ocean. Since whales travel great distances they also cannot be adequately studied in a restricted area; and furthermore, with the new pelagic factory ships the whaling industry was now spreading round most of the Antarctic seas. The work of the *Discovery II* eventually became in effect a survey of the whole Southern Ocean—and this, incidentally, was a task far beyond the capacity of the old *Discovery*. It needed a ship with greater power and cruising range; but since this ocean forms a circumpolar belt, not divided by land masses, the water circulation is relatively straightforward, and in consequence some at least of the problems are simplified. It is worth noting here that the Challenger and other expeditions have given us a broad picture of the oceans as

a whole, and if the work of the *Discovery II* had been planned purely as a further step in the exploration of the oceans—that is to say, with more intensive observations in one ocean—the Southern Ocean might well have been chosen as the region most favourable for the basic study of oceanic phenomena on a large scale.

METHODS AND EQUIPMENT

Ocean surveys of the kind undertaken by the *Discovery II* consist essentially in making observations at suitably scattered points, and interpolating between those points. At intervals varying between about 10 and 200 miles the ship is stopped on 'station', a sounding is taken, and then, with special sampling bottles and deep-sea thermometers, water samples and temperature readings are taken at various levels from the surface to the bottom. Salinity, oxygen content, pH and certain nutrient salts are estimated in the water samples, and density is calculated from the temperature and salinity. At the same time the plankton is collected with vertical and oblique closing nets at different depths, and these sample everything from the shrimp-like 'krill' (*Euphausia superba*), which constitutes the food of whales and many other animals, down to the minute constituents of the phytoplankton (mainly diatoms), which form a basic supply of vegetable food, comparable in a way to the grass on land. From a line of stations these physical, chemical and organic constituents can be mapped in vertical sections of the ocean; a network of stations gives their distribution in three dimensions; and lines repeated at different times of year reveal the seasonal variations. It is a common experience in work of this kind that the more crowded are the observations, the more complex are the conditions found to be. The area of the Southern Ocean is of the order of 30 million square miles, and here, of course, we are dealing with very widely scattered observing points, so that isolated observations, or groups of them, may not be of much significance. However, when the data are studied as a whole an unmistakable pattern does emerge, and certain valid conclusions can be drawn as to the general distribution and movements of the principal water masses, and the way in which they control the distribution of the plankton populations in general and the food of whales, and thus affect the whales themselves.

In deep-sea research we are groping in an invisible world with instruments which can at best give incomplete knowledge, and progress in such investigations is very largely dependent upon, and limited by, progress in the invention of technical devices. Indeed, any new device—echo sounding is an example—will open up new avenues of research, and will often discover phenomena which were not previously known to exist. Even with the methods already available it is obvious enough that for a survey of such a vast area as the Southern Ocean the whole subject of the design and management of ships, and the supply, operation and maintenance of the scientific equipment must loom very large in the organization of the work. Technical problems are greatly increased when oceanic depths are to be explored, and in fact the stage at which any material is examined in the laboratory follows a long series of administrative and practical operations. In this connexion it is significant that the Committee's scientific staff, which usually numbered about eighteen to twenty,

must be supported by an executive staff in the ships of about seventy-five, and a secretarial staff of four or five. I shall explain later how a good deal of the scientific work is farmed out to specialists outside the Committee's organization, but the secretariat also would need to be much larger if we did not have liberal assistance from the common services of a large government department. The management of deep-sea research ships operating far from home involves heavy administrative work, especially before a ship sails on a new commission, but much of this, including arrangements for docking and repairs, placing orders, appointing agents abroad, paying accounts and emoluments, and so forth, has been executed by the Crown Agents for the Colonies, and in our new organization under the Admiralty corresponding services are available.

The *Discovery II* (1036 tons gross) is larger than the average research ship. She is an oil-burning steam ship, with some protection against ice, and has a maximum cruising range of 10,000 miles. The principal items of equipment are the light steam deck engines with reels carrying up to 3500 fathoms of 4 mm. wire for water sampling bottles and vertical nets, and a heavy winch with a small drum for 1000 fathoms of warp, and a large drum for 5000 fathoms of tapered warp. These are for the larger tow nets, trawls, dredges, etc. The instruments and accessories which are used, the sounding machines, the laboratory equipment, etc., could hardly be described here, and it is enough to say that all normal deep-sea operations are provided for, and spares are carried to allow for prolonged and intensive work. The old *Discovery* (736 tons gross) was a wooden auxiliary barque, and her equipment was similar to, though less comprehensive than, that of the *Discovery II*. The *William Scoresby* (324 tons gross) carries a commercial otter trawl and equipment for limited oceanographical work, and she incorporates some of the features of a trawler and some of a whale catcher.

OCEAN SURVEYS BY THE *DISCOVERY II*, 1929-1939

Each commission of the *Discovery II* lasted for about 20 months. In the first of these (1929-31), which was led by Dr Kemp, she worked in the Atlantic sector of the Southern Ocean, but the second commission under Mr John included a complete circumpolar cruise in the winter months, with long V- or W-shaped lines of stations between the southern continents and the ice-edge. The third, fourth and fifth commissions (1933-9), which were led by myself, Dr Deacon and Dr Herdman respectively, included repeated lines of stations on the Greenwich meridian and in 80° W (South-east Pacific), cruises tacking across the summer zone of whales and krill north of the pack-ice, a further circumpolar cruise in summer, and a variety of shorter cruises for subsidiary purposes. Figure 2 shows the ship's principal tracks (omitting some intensive work in the Falklands sector and coastal surveys), and the lines can mostly be taken to represent rows of stations and almost continuous soundings.

The picture of the Southern Ocean which emerges from this work is a series of circumpolar zones or gradients, distorted by the effects of land masses and bottom topography, but recognizable in all longitudes. Any line of stations running from north to south will cut across certain distinct winter masses and several zones with

distinct populations. But from east to west there is little change, and the major features of the Southern Ocean can be seen in a vertical section in almost any meridian.

It was already known that the general drift in this ocean is from west to east (except for a counter current close to the Antarctic coast), and the German research ship *Meteor* in 1926 showed that in the Atlantic sector there is an Antarctic 'con-

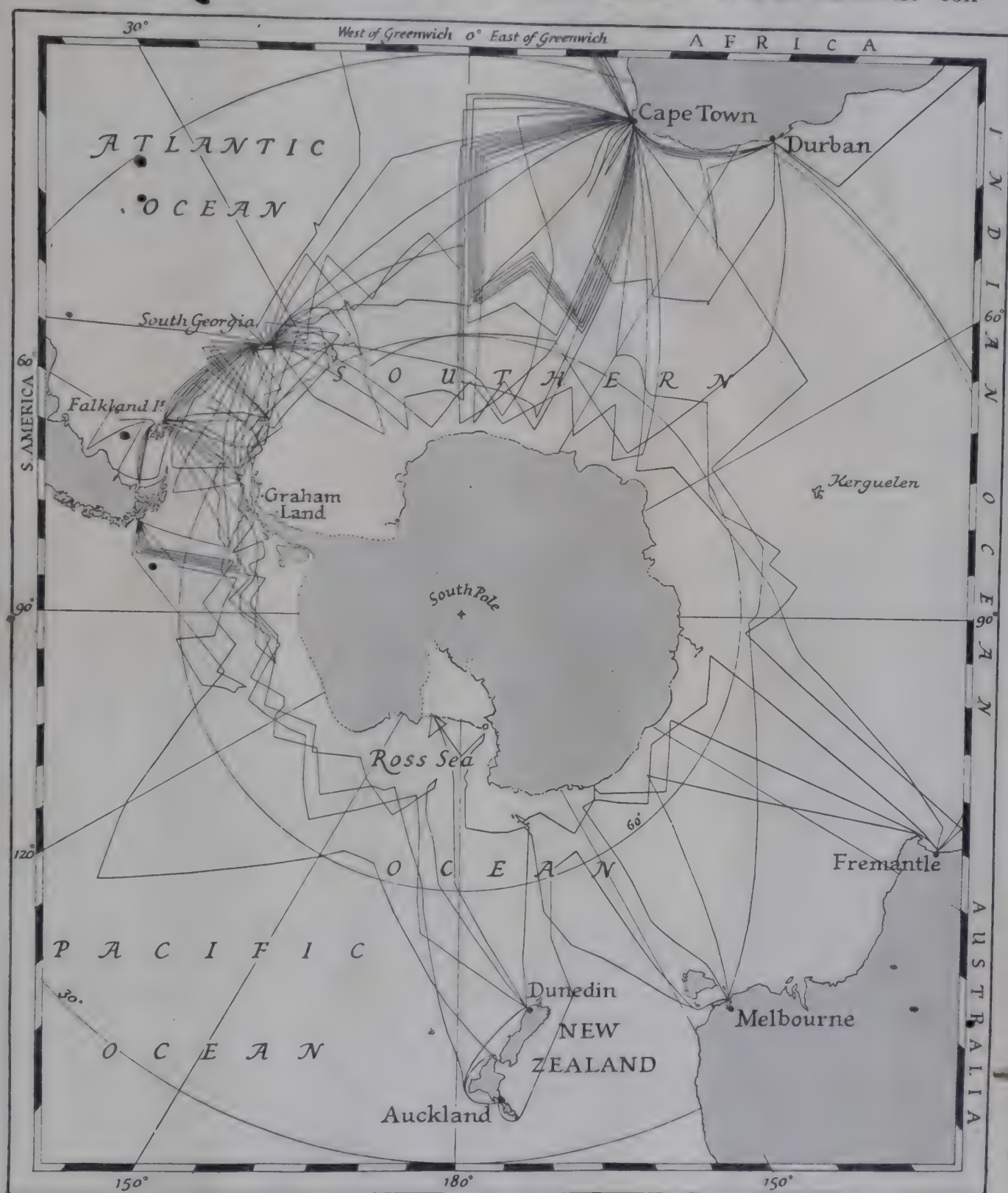


FIGURE 2. Principal voyages of the *Discovery II*, 1930-9. (The tracks are only approximately correct. Where they are crowded together they have been straightened and separated.)



FIGURE 4. The Marine Station at Grytviken, South Georgia.



FIGURE 5. The Laboratory at the Marine Station.



FIGURE 11. Biological laboratory in the *Discovery II*.

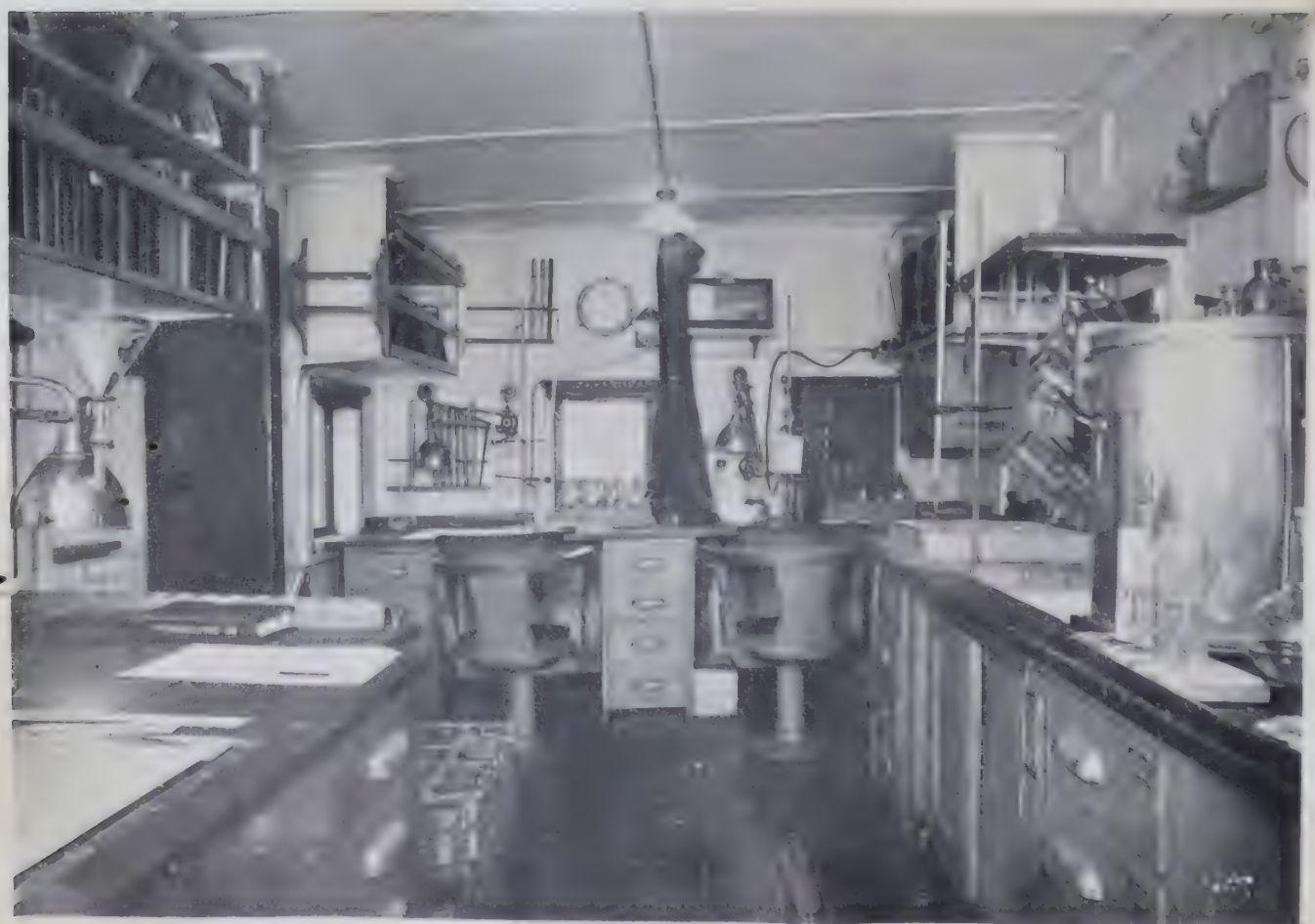


FIGURE 12. Chemical laboratory in the *Discovery II*.

downwards to the north, underneath the sub-Antarctic surface water. Hence it is inferred that the position of the Antarctic convergence at the surface is determined by relative movements of the intermediate and bottom currents.

In the upper part of the Antarctic surface layer there is a rich growth of diatoms and animal plankton. Organic matter sinks and enriches the deep intermediate water with nutrient salts which are carried south and back to the surface. The key to the distribution of pelagic organisms in the Antarctic, including the food of whales, has been found in the relative movements of the surface and deep intermediate layers. We have found that the various species of the animal plankton have different devices for spending part of their lives moving east and north in the surface layer, and part moving east and south in the deeper water. It can scarcely be doubted that this is the mechanism by which they maintain the boundaries of their normal distribution—some by diurnal migrations from one layer to the other, many by spending the summer in the surface and the winter at a deeper level, and some spending different parts of their life cycle in the two layers. Apart from the shoaling krill (*Euphausia superba*), those species which form the bulk at least of the macroplankton exhibit the second of these methods. Concentrated in the surface layer in summer they almost desert it in winter, having sunk far into the southward-moving intermediate water; and both the sinking in autumn and the upward movement in spring take place earlier in the lower than in the higher latitudes. This seasonal migration must set up a large-scale circulation which brings about the replenishment of the stocks in the far south. The krill gives an interesting example of the third method. The adults live very close to the surface, but their eggs sink to 1500 m. or more, and the successive larval stages are passed between increasingly shallow horizons as they migrate through the intermediate layer towards the surface. It must be supposed that in the general eastward drift, each species, in a succession of generations, will circulate perpetually around the Southern Ocean, and that by adjustment of its movements between the two layers it can keep within a more or less restricted zone with a suitable range of temperature and perhaps other conditions. This may require final confirmation, or it may be an over-simplification of what actually takes place, but it does seem to offer an adequate explanation of the principal features of the plankton distribution.

DIRECT RESEARCH ON WHALES

Turning now to the more direct investigation on whales I should say something of the method of marking them. This was an entirely new undertaking, but after considerable experiment had been made it was found that a stainless steel tube fired from a shot gun was successful. The marks are designed to bury themselves under the blubber, and each has a serial number and an inscription offering a reward for its recovery by the whalers. About 5000 whales have been marked in various parts of the Southern Ocean, mostly by the *William Scoresby*, but also by hired whale catchers near South Georgia, and so far nearly 300 marks have been recovered. (An unknown proportion may be lost through a healing process.) These recoveries have had several important results. They have clearly proved that at least some whales do (as was believed) undertake long migrations between the Antarctic and tropical

waters; they have shown that a whale usually, though not always, returns after its winter migration to the same part of the Antarctic in the following summer; and they have shown that the species differ in their tendencies to become segregated into distinct communities. The proportions in which the marks are recovered from the different species are instructive; and as time goes on the marks will no doubt provide information on the ages and length of life of whales. Some marks recovered recently had been carried by whales for periods up to 14 years.

The anatomical examination of whales at South Georgia was extended to South African whaling stations for two winter seasons, and more recently to the modern pelagic factory ships, and altogether some 6000 whales have been measured and examined. The majority of these were fin and blue whales, but humpback, sperm, sei, and a few right whales were included. The blue whale (*Balaenoptera musculus*, figure 7, plate 2) is the largest species, with a maximum length of about 100 ft., and the fin whale (*B. physalus*)—not very much smaller—is the most plentiful. These two, and in some regions the humpback and sperm whales, are the mainstay of the whaling industry. From this work a large body of statistical data has been obtained on the range of variation within the species, and information has been gathered on their food, parasites, etc., but attention has been paid more especially to their breeding, growth and age, and to the constitution of the local populations of whales. The breeding season, period of gestation, rate of breeding, and average lengths at birth, sexual maturity and full growth have been approximately established, and provisional estimates of the rate of growth have been made. As a general rule they appear to grow fast but to breed rather slowly. Methods have been found also for judging the relative, though not yet the absolute, age of individual whales.

The method of age determination will serve as an example of the technique of work in whaling factories. The best criterion we have found is the number of old corpora lutea, or corpora albicantia, in the ovaries of an adult female.* Each of these corpora lutea records the shedding of an egg from the ovaries, and there is good evidence that they persist for many years if not through the life of the whale. The number of them in a pair of ovaries can be correlated with the attainment of full growth (which can be checked by examination of the vertebral epiphyses), and so we can take this number as at least a rough indication of the relative age of a female whale. We cannot yet translate numbers of old corpora lutea with certainty into years of age, but the average rate of accumulation is probably about one a year. The physical work involves excavating the genitalia from a huge mass of viscera, and cutting notches in the vertebrae with an axe.

The work in whaling factories is carried out in unattractive conditions, but it does not require the elaborate equipment and technical experience needed for the deep-sea oceanographical work. The reason is of course that we have not had to catch our whales. This has already been done for us by the whalers, and I must acknowledge our great indebtedness to the whaling companies (too many to name here) who have not only given us many facilities for our work, but have regularly collected specimens for us over a series of years.

* Professor Ruud, of Oslo, has recently found a method of estimating the age of young whales by examination of the baleen plates.

From the data collected in whaling factories we could learn much more than we have learned about the constitution of the populations of whales if the catch were a fair sample of the whales in the sea. The whalers take whales where they can find them, and in a broad sense the catch is taken at random. But to some extent it must be biased by varying economic and political factors, weather conditions, and the whaling regulations, and it is often difficult to know how much, and what form of, selection may result. For example, there was a gradual increase over a series of years in the percentage of pregnant blue whales in the Antarctic catches. It was difficult to find any explanation of this in changes in the areas, methods, or times of hunting, and it almost seemed that the rate of breeding was actually increasing; and yet the explanation might lie in some small and gradual change in the method of hunting—perhaps a tendency for the factories to work, say, a little farther away from the pack-ice edge, and so to sample a slightly differently constituted population. However, there are conclusions about the populations which can be drawn with some confidence, and we have good evidence on such matters as the relative abundance of the different species in the Antarctic population, of the sex ratio in baleen whales, and of certain changes in the make-up of the population during the summer season.

The Discovery Committee's work on whales has generally helped to give a solid foundation for the regulation of whaling, and many of the conclusions reached have specifically affected the regulations or assisted in the interpretation of the statistics of catches. For instance, simple estimates of the average length at sexual maturity have contributed to the protection of immature whales, and allowed the percentage of immature whales in the whole catch to be calculated from year to year. Observations by both ships, and the recoveries of whale marks, have provided evidence not only that blue whales are much scarcer than fin whales, but that a larger proportion of the stock of blue than of fin whales is killed. Counts of old corpora lutea have shown significant changes in the age distribution of the catches of blue whales, and it is partly in consequence of such findings as these that the protection of this species has been strengthened. Directly or indirectly the work has further contributed to the delineation of sanctuaries, an overall limit to the total catch, measures for the protection also of humpbacks, and the limits to the open season. The international regulations are not of course based only on the Discovery Investigations. The Norwegian scientists for instance have made important contributions, especially in the analyses of the statistics of catches.

OCEANOGRAPHY AND THE DISTRIBUTION OF WHALES

Although the survey of the Southern Ocean has become, in its own right, one of the principal aspects of the whole programme, the oceanographical work is an essential part of the research on whales. Without a knowledge of the oceans in which the whales live we could not see our problems in their right proportions, but more specifically this knowledge is concerned with their distribution, movements and habits. The majority of whales spend the summer in polar waters where they concentrate on the rich feeding grounds, and in winter there is a migration to lower

latitudes where breeding takes place. It is assumed that they seek warmer water for breeding, and that finding little to eat, they depend for their energy on the fat stored in their blubber and other tissues. In the summer their distribution must depend largely on their food, though they are directly influenced also by such factors as the position of the pack-ice and almost certainly the sea temperature. Cause and effect are to be examined, and it is important to have one's objects clearly in mind. The ultimate aim, I should say, is prediction. We should like to be able to say where, when, and in what quantities whales will be found, what will be the effect of hunting them at one time and place or another, and so on. For this we must try to relate the pattern of distribution of the plankton and whales on the one hand to the more readily defined physical and chemical pattern of the sea on the other, and to connect the fluctuations in their occurrence with fluctuations in such conditions as the ocean temperatures and currents. We should even try to go further back to the primary factors—the climate, and the winds and other variable weather conditions which affect the sea. This is a long road to travel, and the physics and chemistry of the Southern Ocean still need further study. However, some definite progress has been made. Knowledge of the plankton circulation is a step forward, and we are in a position now to draw fairly reliable maps of the distribution and density of the krill and the phytoplankton in the Southern Ocean, and of the principal physical and chemical conditions on which their distribution depends. Indeed, we could now, up to a point, define in oceanographical terms the circum-polar zone which whales seek in the summer months. This can be better done when certain gaps in the general survey are filled, but in many ways we can go further than this. Some years ago, for instance, Professor Hardy showed how the local distribution of whales could be predicted from the distribution of phosphates in the water. A number of other correlations have been explored, and although there is still much to be done I think we shall be able to link up more and more the oceanographical work and the direct observations on whales.

SPECIAL AND SUBSIDIARY INVESTIGATIONS

The Discovery Investigations have been focused very largely on the biology, or perhaps I should say the bionomics, of whales, and the ocean surveys might never have been undertaken were it not for their bearing on the organic resources of the Antarctic seas. Although there has been no departure from the original objects, the work has in many ways spread far beyond them, and its ramifications have extended into many other fields. This I think commonly happens when a problem is approached on a broad basis, and certainly where well-equipped research ships are concerned. The principal work of the *Discovery II* has been to survey the major features of the Southern Ocean; this in itself is a step in the general exploration of the sea, but it is a type of survey which can also take in its stride many special and subsidiary investigations. Much new information, for example, has accrued on the distribution and seasonal changes in the Antarctic sea ice, innumerable deep-sea soundings have revealed many varieties of bottom configuration; large collections of the marine fauna and flora have been added to the national collections at home;

the biology of seals has been studied; coastal surveys and soundings have improved the Admiralty charts; and the latest edition of the *Antarctic Pilot* is full of contributions from the officers of the research ships. Material obtained for others to work on includes routine observations for the Meteorological Office, and the collection of specimens for teaching departments and individual research workers. The *William Scoresby* also, in addition to her principal work of marking whales and assistance in general oceanography, has undertaken two major operations: an exploratory trawling survey, divided into three periods, of about 150,000 square miles of the Patagonian continental shelf, and an oceanographical survey of the Peru coastal current. Marketable fish were found during the trawling, but not in sufficient numbers to offer very good prospects of a commercial fishery. The value of the Peru current survey lay mainly in the accurate delineation of the current and in the measurement of its physical features, especially of the subsurface layers about which little was previously known.

ADMINISTRATION AND RESEARCH AT HOME

In giving a brief account now of the Committee's organization at home I must speak in the past tense, for although the scientific work both at home and in the field is continuing on much the same lines as before, the administrative system under the Admiralty is of course not quite the same as it was under the Colonial Office. The Discovery Committee was constituted as a team which included members qualified to advise not only on purely scientific matters but also on such subjects as the whaling industry, navigation and the management of ships, the Antarctic regions and organization of expeditions, and finance; and each member played an active part in the Committee's work which sometimes involved rather more than attendance at meetings. From 1937 onwards the meetings were also attended by observers representing the governments of Canada, Australia, New Zealand and South Africa. Since the scientific programmes broke new ground in many directions as time went on, and the organization of what might be regarded as a series of expeditions involved much administrative work, it was found convenient to hold rather frequent meetings of the main Committee (about nine times a year) to decide matters of principle and to receive the recommendations of a Scientific Subcommittee and sometimes of a Ship Subcommittee.

The headquarters of the Discovery Investigations have always been in London, and although we have not possessed premises in which all the administrative and scientific work could be brought together, a system was evolved which worked very smoothly. Committee meetings were held in the Colonial Office, the scientific staff were accommodated in the British Museum (Natural History), and the Director and Secretary had offices in Queen Anne's Chambers, in Westminster, within very easy reach of the Museum on the one hand, and the Colonial Office, Admiralty, Fisheries Department, etc., on the other.

The results of the Committee's work are published in the *Discovery Reports*, which are printed by the Cambridge University Press. There has been a more or less steady output of these reports since 1929, and so far twenty-five volumes have been issued.

At sea, or elsewhere in the field, the original data receive some preliminary treatment, but this of course must be followed up by prolonged study of the material at home before papers are ready for publication. A ship like the *Discovery II* can accumulate data at an embarrassing rate, but although there is still a considerable mass of material which has not been examined, we have been able to keep pace with the more important work and to deal gradually with the less urgent aspects of it. This however could not have been done without the generous co-operation of other bodies and persons, and we are specially indebted to the Trustees and staff of the British Museum (Natural History) who have provided accommodation for our staff and collections, and given constant advice and help in our work.

Most of the research on the biology of whales, and on the plankton and physical oceanography of the Southern Ocean—that is to say in the main stream of the Discovery Investigations—has been carried out by members of the Committee's scientific staff (or sometimes by former members). The duties of the staff have varied from time to time, and each member has had experience of several branches of the work. Of about fourteen scientific officers in the years before the war the majority were zoologists and three were chemists; and most of their time in London was free for research, though there was also some work to be done in the supervision of the collections, library and scientific equipment and occasional assistance in the Director's office. The *Discovery II* generally carried two or three zoologists, and one or two chemists, with three assistants and certain members of the crew detailed for scientific duties. In the *William Scoresby* there might be one or two zoologists, and others might be working in whaling factories or on other detached missions. At any one time, however, about two-thirds of the staff would be working in London, or more when the *Discovery II* was home. Naturally we do not attempt to work on the whole of the data from one voyage before dealing with material from the next voyage; the method is for any one worker to take up some particular subject at a suitable stage (say the seasonal periodicity in the phytoplankton, or the bottom topography of some region) and to work through all the relevant data available up to date. The choice of such subjects depends partly on their importance in the co-ordinated plan of research, partly on the interest in or flair for any particular line of work which a member of the staff may develop, and partly of course on the time when sufficient data have been collected.

Although most of the scientific staff have spent considerably more time at home than at sea, it would not be possible for them to deal with every aspect of the work; but we have been greatly assisted by specialists in various institutions here and abroad, or working independently, to whom we have been able to farm out much of the material. Of these about forty have so far contributed one or more monographs to the *Discovery Reports*, and a number of others are examining various parts of the general collections. Many of these contributions are systematic reports (prepared in the style of the *Challenger Reports*) on various groups of marine animals; and although the material on which they are based was usually obtained with little interruption of the field programmes, they have added very considerably to what is known of the marine fauna. Other contributions from specialists include papers on such subjects as bottom sediments, special studies of invertebrates, the geology of

the Falkland Islands Dependencies, and so forth. All papers are published in the order in which they are received, and are distributed to many libraries at home and abroad, the balance being kept for sale and incidental presentations.

The work of the Discovery Committee is to be judged mainly by its contribution to oceanography and to the biology of whales, but the Committee has also had advisory functions. They have always been represented at international conferences on whaling, and have been able to give advice and assistance, not easily supplied from other sources, to other expeditions, departments and institutions.

THE PRESENT AND FUTURE

To review in a few words the ground which has been covered up to the present by the Discovery Investigations, I should say that four main points might be made. First, a survey—not quite complete yet—has been made of a large ocean area and the seasonal changes which take place in it. Secondly, I think a real step forward has been made in relating the oceanic fauna to the currents and water masses. Thirdly, researches on the biology of whales have helped the international regulation of whaling. And fourthly, a pool of information on the Southern Ocean and Antarctic regions has been built up.

I have emphasized the intermediate position of the Discovery Committee's work between an expedition and an institution, but it might be more correct to say that it has been in a state of transition, for it is now approaching the status of a permanent service. Since 1939 until recently the ships were on charter to the Ministry of Transport, but the work at home has been continued with a managing committee under the chairmanship of Mr J. M. Wordie, one of the original members of the Discovery Committee. In the National Institute of Oceanography, to which Dr Deacon has been appointed as director, we have now joined a larger organization which is to 'advance the science of oceanography in all its aspects', and as a going concern we can give it a flying start. This does not mean that the Discovery Investigations will lose their identity in the new Institute; in fact, the title will be retained at least for the continuation of the work in the south with which it is associated, and the *Discovery Reports* will continue to be issued. Under a provisional executive committee, acting for the Institute, our affairs, indeed, are moving fast. Additions have been made to the scientific staff, the *William Scoresby* is commissioned again, and the *Discovery II* will soon be ready for sea. During the first year or two the work of both ships will be in continuation of the Discovery Committee's programme, though some preliminary work in a wider field will be included. Thus the principal task of the *Discovery II* will be to round off the general survey of the Southern Ocean, but work will also be started in the central Indian Ocean. The marking of humpback whales during their winter migration off the north-west Australian coast will be the main item in the *William Scoresby's* programme, but she is starting with a preliminary survey of the Benguela current, and will test some possible methods of catching the famous Coelacanth fish, *Latimeria*, off East London. Although the ships are occupying most of our attention just now, we have in mind plenty of new developments in the more direct observations on whales.

Although much of the data to be collected in the near future must be properly comparable to what we have obtained in earlier years, we shall make use of some of the new instruments which have been developed more recently. The bathythermograph, for instance, can hardly be omitted now from the outfit of a research ship, and the *Discovery II* will also carry the Kullenberg piston core sampler and instruments for seismic sounding in bottom sediments, both of which offer many possibilities.

Deep-sea oceanography has travelled a long way since the time of the Challenger Expedition, and it seems inevitable now that a ship should specialize to some extent, at least within any given period or voyage. In our experience fruitful results are obtained from planned voyages in which the ship's principal movements are dictated by some major object, but which form a framework within which some attention can be paid to many more specialized problems. The major object might, of course, be a general survey of some ocean area, or it might concern a particular subject such as the exploration of the sea floor. The Discovery Investigations will in the future no doubt take their place along with other long-term researches in the physical and biological aspects of oceanography undertaken by the National Institute. We have been concerned over the years with a wide field of inquiry, ranging from problems in pure physical oceanography to such practical questions as the machinery of the regulation of whaling, and we have a flexible organization. It has all been a great adventure in the past, and I believe that it will be so in the future.

Studies on the ionization produced by metallic salts in flames

II. Reactions governed by ionic equilibria in coal-gas/air flames containing alkali metal salts

BY H. BELCHER AND T. M. SUGDEN

Department of Physical Chemistry, University of Cambridge

(Communicated by R. G. W. Norrish, F.R.S.—Received 6 October 1949)

The ionization produced by adding alkali metal salts to coal-gas/air flames has been studied by measuring the attenuation of 3 cm. waves, this attenuation arising from free electrons. In a previous paper (Belcher & Sugden 1950) the coefficient relating attenuation to electron population has been determined and is used in applying the present results to the problems of chemical equilibrium involving ions in the flames.

An examination has been made of the degree of equilibrium obtained in a number of chemical reactions in the flame by varying the electron content, the amount of added alkali metal salt, and its nature. Flame temperatures were measured by the sodium line-reversal method, considered to be the most suitable for the purpose. In the case of potassium bicarbonate the results are compatible with chemical equilibrium of the ions if suitable account is taken of other equilibrated reactions involving potassium hydroxide, oxide and hydroxyl ions.

Some evidence of disequilibrium has been observed when potassium halides are added, where both formation of negative halide ion and halogen acid affect the electron content. This appears to be connected with concentrations of hydrogen atoms and hydroxyl groups (chain centres) which diverge from the values expected at equilibrium. Addition of extra free halogen confirms the anomalies. In the case of the chlorides and fluorides, the halogen appears to be completely removed as halogen acid.

A comparison of the alkali chlorides on the whole supports the use of the Saha equation relating equilibrium ionization and flame temperature, providing account is taken of side reactions.

1. GENERAL AND INTRODUCTORY

In a previous paper (Belcher & Sugden 1950) it has been shown that the electron concentration in flames may be measured very conveniently by the attenuation of electromagnetic radiation of 3 cm. wave-length, and is given by the equation

$$n = 2.6 \times 10^{11} \beta,$$

where n is the number of electrons per cm.³, and β is the attenuation expressed in decibels per cm. thickness of flame, provided that the attenuation does not exceed 2 db./cm. Previous work (see Wilson 1944) on d.c. conductivity of flames has tended to show that when an alkali metal salt is added to a coal-gas/air flame, the ionization may be represented by the Saha equation (Saha 1920), which is a thermodynamic relation relating ionic equilibrium with temperature:

$$\log_{10} K_p = -5050 \frac{V}{T} + \frac{5}{2} \log_{10} T - 6.9,$$

where V is the ionization potential of the alkali metal A , T the absolute temperature ($^{\circ}$ K) and K_p the equilibrium constant of the system $A \rightleftharpoons A^+ + e$, the concentrations being expressed in atmospheres of partial pressure. If this is the only important ionic equilibrium in the flame, then $[A^+] = [e]$, and if the ionization is small, then $[A] \sim [X]$, where X is the total salt added, expressed as a partial pressure in atmo-

spheres, and the constant $K_p = [\epsilon]^2/[X]$. From measurements of $[\epsilon]$ it should be possible to obtain V for the metal knowing the flame temperature, or otherwise a value of T for a flame should be determinable from the Saha equation, given the ionization potentials of the alkali metals.

The present work was undertaken to determine whether this simple system was sufficient to explain the observed behaviour in flames, and if not, whether and to what extent it was modified by side reactions of A , A^+ and ϵ , and if the postulate of temperature equilibrium for all the reacting bodies was a valid one. The region of flame studied for premixed coal-gas/air flames was the extended zone immediately above the small blue cones, which make up the primary zone of reaction in the flames. In the narrow blue zone of reaction, nothing approaching temperature equilibration is expected, since it will have no clearly defined temperature, but in the extended zone above, almost colourless in the absence of added salt, some degree of equilibrium may be expected, and it is with this zone of hot gases that the present ionization work is concerned. A detailed description of the flame and the measurement of attenuation has already been given (Belcher & Sugden 1950). A mixture of coal gas and air in known proportions, to which a fine spray of alkali metal salt was added from an atomizer, was burned in a Meker type burner, giving a flame 18 cm. long and 1 to 2 cm. thick, and the attenuation of 3 cm. waves by the added salt measured. No measurable attenuation was observed in the absence of added salt.

For a given flame, $[\epsilon]$ is proportional to the attenuation β db./cm., and $[X]$ is proportional to the concentration of the salt solution in the atomizer, so if the Saha equation alone is important, then a plot of $\log \beta$ against $\log N$ (normality of solution), should be a straight line of slope 0.5. It became obvious early in the work that this was not so in the flames studied (see, for example, figure 3).

In order to be able to place this and other divergences on a quantitative basis, other measurements in the flame were necessary, viz. those of $[X]$, the concentration of salt added and T , the temperature of the flame. The methods used for these quantities are described below. In the light of these measurements and the attenuation results, an attempt is then made to estimate the extent to which various side reactions affect the electron concentration for salts, which, in the case of the halides, include the tendency of the halogen to form negative ions and the effect of the hydrogen in the flame in removing the halogen as the acid, and even in the absence of halogen, the formation of alkali hydroxide and oxide and negative hydroxyl ions in the flame. When sufficient account of these has been taken, the results are then examined to find whether it is possible to interpret them as consistent with temperature equilibrium of all the reactions involved, and if not, in what reactions disequilibrium occurs. In this way it is hoped that ionization measurements may be able to throw light on some of the chemical reactions taking place in the flames.

2. DETERMINATION OF THE SALT CONCENTRATION $[X]$

A determination of the total salt concentration added to the flame, expressed in atmospheres $[X]$, was made by a photoelectric method. The flame was viewed through a rectangular tube about 2 ft. long and $1 \times \frac{1}{2}$ in. in cross-section arranged perpendicular to the flame surface (figure 1).

Light passing down the tube (matt black inside) fell on a photocell with a cathode much larger than the end of the tube. The cathode of this cell was connected to the grid of a triode and to one plate of a condenser, so that the total light falling on the cathode determined the potential of the grid. The triode was in a cathode follower circuit and calibrated in terms of integrated light falling on the photocell cathode against the current passed by the valve. Since the photocell had a filter in front of it which only allowed the sodium D lines to pass out of the total radiation present, it thus acted as an integrator for this radiation. Small weighed beads of fused sodium chloride were introduced into the flame as in figure 1, so as to give a yellow streak which filled the whole thickness of the flame at the level of observation, the streak at the same time being narrower than the measuring tube. The total light obtained from beads of different weights gave an accurately linear calibration for any given

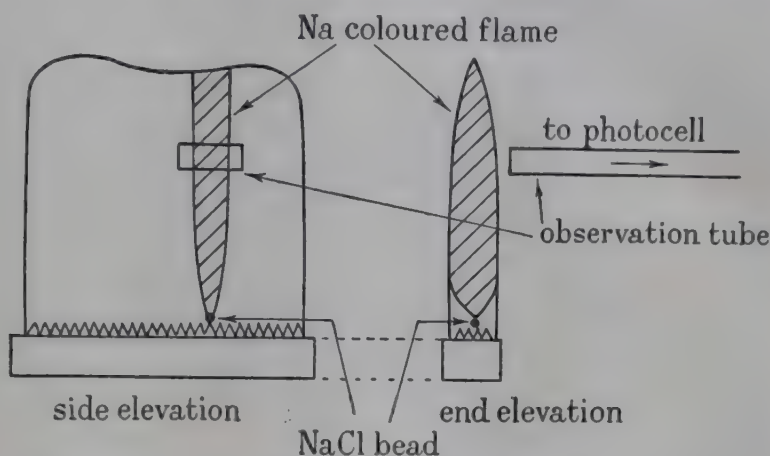


FIGURE 1. Arrangement for determination of $[X]$.

flame. A solution of known strength of sodium chloride was placed in the atomizer, and the total sodium D radiation collected by the tube measured in a given time (10 to 20 sec.). It was thus possible to say what weight of sodium per second crossed the width of the tube, and hence the weight crossing any centimetre width of the flame.

In order to obtain the instantaneous concentration $[X]$ it is necessary to know in addition the vertical velocity of the flame gases, and the thickness of the flame at its base. The latter was 1.6 cm., and the former was measured by a method described by Lewis & von Elbe (1943). A small quantity of fine aluminium powder was introduced into the gases before they reached the burner, giving luminous streaks in the flame, which were photographed with a camera with a calibrated shutter. Hence the velocity of flow could be obtained. The weight of sodium chloride introduced into each cubic centimetre of flame gases could thus be computed and converted by the appropriate factors to a partial pressure in atmospheres. For this conversion the reversal temperatures were used (see below). The atomizer was set so that for normal solution $[X] = 3.2 \times 10^{-5}$ atm. This figure has a fair probable error (of the order of 10 %) on account of the large number of measurements involved in its determination. However, it is slightly lower than values obtained on the basis of atomizer loss. Previous workers have used this loss as a measure, but it was considered advisable to check this by another method since some spray, although not a great deal, was

deposited on the sides of the tubes. Bennett (1927) used a method to obtain the velocity in which the distance between timed puffs of sodium chloride into the flame were measured, but such a method was found to be much less accurate than the use of aluminium powder.

3. PREDICTION AND MEASUREMENT OF FLAME TEMPERATURE

The idea of flame temperature is one whose connotation still remains not entirely clear. In the narrow zone of primary reaction there is no definite temperature since there is no thermodynamic equilibrium, but in the extended zone of hot gases above this, which is studied in the present work, some degree of equilibrium is expected, and hence it becomes possible to speak of its temperature, although it is by no means certain that complete equilibrium between all the degrees of freedom of the flame gases occurs.

If the composition of the unburnt flame gases, the heats of formation of the initial and final substances, and the specific heats of the products from room temperature to a suitably high temperature are known, it is possible to compute, by standard thermodynamic methods, a value of temperature which would hold good if complete equilibrium obtained together with no heat loss by entrainment of air or radiation. This has been done for the range of coal-gas/air mixtures studied, assuming the air to be saturated with water vapour at 20°C, leading to the upper line ($T_{\text{calc.}}$) of figure 2. The principal constituents are expressible by the following equilibria (no suffixes are used for K 's):

$$K = \frac{[\text{CO}][\text{H}_2\text{O}]}{[\text{CO}_2][\text{H}_2]}, \quad K = \frac{[\text{OH}]\sqrt{[\text{H}_2]}}{[\text{H}_2\text{O}]}, \quad K = \frac{[\text{O}_2][\text{H}_2]^2}{[\text{H}_2\text{O}]^2}, \quad K = \frac{[\text{H}]^2}{[\text{H}_2]},$$

the amounts of nitric oxide and atomic oxygen and nitrogen being negligible in these flames. It is very important in addition to have some measured value of the flame temperature.

A choice of methods is available for this, from among which it was decided that the sodium-line reversal method was the most suitable for various reasons. The method has been described frequently (see, for example, Griffiths & Awberry 1929), is simple and reproducible, and has the advantage of not requiring any solid object to be placed in the flame. It is also a relevant method since it uses an alkali metal as its basic substance, as was also used in obtaining the electrons for attenuation measurements. It relies on the supposition that thermal equilibrium is reached between excited and unexcited sodium atoms in the flame, and hence that the spectral brightness of the sodium D lines will be proportional to that of a black body at the temperature of the flame gases. If these conditions are not obeyed, the results will be erroneous, or even meaningless. Previous workers have shown (see Lewis & von Elbe 1938) that the method gives results which are somewhat below the calculated values, the divergence being greatest in the region of stoichiometric compositions: this may arise from entrainment of air, or possibly as David (1935) and his co-workers have suggested, from retainment of considerable energy in 'latent' form—i.e. in degrees of freedom out of equilibrium with the rest. At extreme compositions

Jones, Lewis, Friauf & Perrott (1931) have found results higher than the calculated ones, and it is difficult to reconcile this with complete equilibrium of the sodium. Nevertheless, this method appears to be the most reliable for determining this rather ill-defined quantity, and some degree of correlation between temperatures and attenuation might be expected.

A slight variant of the usual technique was adopted. An enlarged image of the filament of a tungsten 'ribbon' lamp (6 V, 108 W) was formed at the height of the wave-guides and across the thickness of the flame, so that the region of it corresponding with the central ribbon extended across the central two-thirds of the flame.

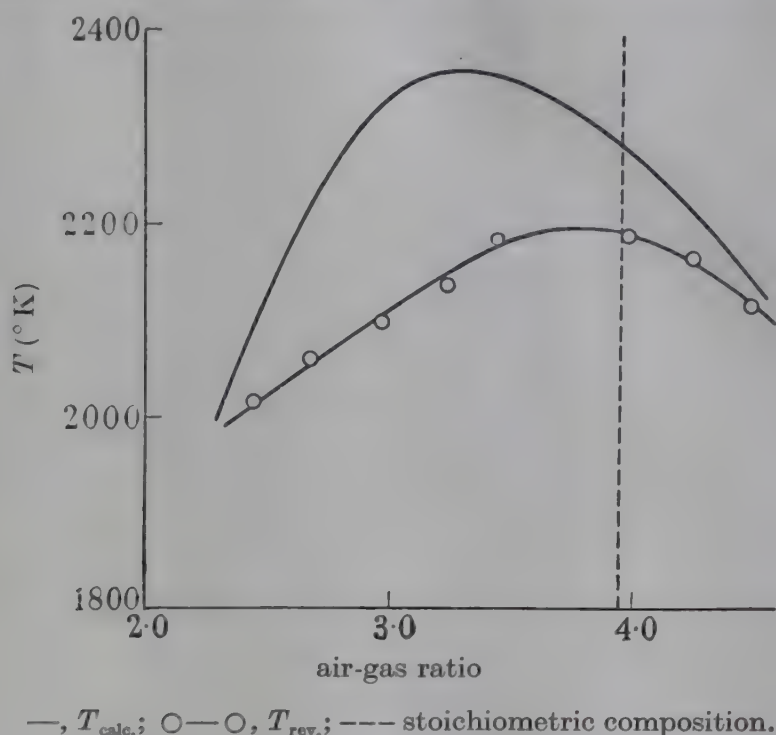


FIGURE 2. Flame-gas temperatures.

This was viewed by a spectroscope via a slit arranged in such a manner that only this portion of the image gave light falling on the spectroscope slit, and also so that the only light from the flame reaching the spectroscope was contained in the solid angle subtended at the spectroscope slit by the specified portion of the image. In this way a suitable average over the region of flame studied in attenuation was obtained. The brightness of the lamp was adjusted until the sodium D lines (from sodium chloride solution in the atomizer) were just invisible against the continuous background. The mean position of disappearance of the lines was taken from readings approaching from either side both visually and photographically. Accuracy was $\pm 5^\circ \text{C}$. The ribbon lamp was calibrated against a standardized disappearing filament pyrometer, and a correction applied for the emissivity of tungsten at the sodium D wave-length and the pyrometer calibration wave-length, and also for the glass lens used in the flame experiments.

The results are plotted as $T_{\text{rev.}}$ in figure 2 and show the usual kind of relation with the theoretical curve. These figures will be taken as a basis in the first place for considering flame equilibria. In table 1 are given values of $[\text{OH}]$, $[\text{O}_2]$ and $[\text{H}]$ as calculated for the flames at $T_{\text{calc.}}$, and also at $T_{\text{rev.}}$, assuming the calculated mixture to

TABLE 1. COMPOSITION OF FLAME GASES IN OH, O₂, H

(All concentrations expressed in atmospheres.)

air/gas ratio	[OH]		[O ₂]		[H]	
	$\times 10^3$ at $T_{\text{rev.}}$	$\times 10^3$ at $T_{\text{calc.}}$	$T_{\text{rev.}}$	$T_{\text{calc.}}$	$\times 10^4$ at $T_{\text{rev.}}$	$\times 10^3$ at $T_{\text{calc.}}$
2.44	0.7	1.0	1.2×10^{-6}	1.5×10^{-5}	4.2	1.1
2.67	1.8	4.2	6.7×10^{-6}	7.0×10^{-5}	4.7	1.5
2.93	3.6	11.0	3.3×10^{-5}	3.0×10^{-4}	4.9	1.6
3.23	5.4	14.0	8.3×10^{-5}	7.5×10^{-4}	5.1	1.6
3.58	7.7	14.4	2.6×10^{-4}	1.5×10^{-3}	5.2	1.4
4.00	9.0	12.4	7.8×10^{-4}	2.0×10^{-3}	5.0	1.1
4.25	8.9	11.2	2.0×10^{-3}	2.3×10^{-3}	4.2	0.9
4.50	8.8	10.4	5.0×10^{-3}	6.0×10^{-3}	3.0	0.7

be cooled to $T_{\text{rev.}}$ from $T_{\text{calc.}}$. These figures will be found to have importance later when considering the attenuation results. They are likely to be in considerable error, being minor constituents, but will show all the trends which actually occur in the flames. The thickness of the flames was measured by placing a platinum wire across them at the height of the wave-guides and measuring the length of it which became white-hot. The boundaries of this region were quite sharp and readings were reproducible to 0.05 cm. It is considered that this is a suitable measurement for an ill-defined quantity such as flame thickness, since it measures the thickness of the hot parts of the flame—in which most of the ionization occurs.

4. A STUDY OF POTASSIUM SALTS

(a) Results

Various types of flame system have been studied, the results of which will be described with reference to particular variables. As a basic reference substance, potassium bicarbonate was used in the atomizer, since it does not give rise to any anionic substance which will react with electrons in the flame, except possibly hydroxyl, already present in considerable quantity in the flame gases. The results it gives are plotted in figure 3 as a series of graphs of the attenuation (β db./cm.) against the normality (N) of the solution in the atomizer (logarithms in both cases).

β is directly proportional to n , the number of electrons per cm.³ of flame, and $[X]$, the concentration of the salt, expressed as a pressure in the flame, is proportional to N . The proportionality constant for conversion of β to n will not vary significantly over the range of flame compositions, as is indicated by the constant ratio of attenuations at different wave-lengths, and as would be expected from the slight variation in mean collision diameter over the composition range, but in order to express n as $[\epsilon]$, the electron pressure in atmospheres, a temperature correction will be necessary in comparing results at different flames, whose reversal temperatures ranged from 2000 to 2200° K.

It will be seen that at low concentrations the results show a limiting slope of 0.5 in figure 3, which is in agreement with the Saha equation, expressed as $[\epsilon]^2 = K_1[X]$. A quantitative test may be made by comparing $[\epsilon]^2$ as measured and as predicted

from values of $\dot{K}_1$, calculated from the full Saha expression, and in the linear regions of figure 3 it is found that the measured $[\epsilon]^2$ is always about one-fifth of the predicted one. Figure 4 has been plotted for $2 \log_{10} [\epsilon]$ against $1/T$ for the dilute solution case ($N/8$), where T is the reversal temperature, and gives a reasonably straight line with a slope corresponding with an ionization potential of 4.0 ± 0.2 eV (full line), exemplifying the following form of the Saha equation:

$$2 \log_{10} [\epsilon] = -5050 \frac{V}{T} + \frac{5}{2} \log_{10} T - \text{constant}.$$

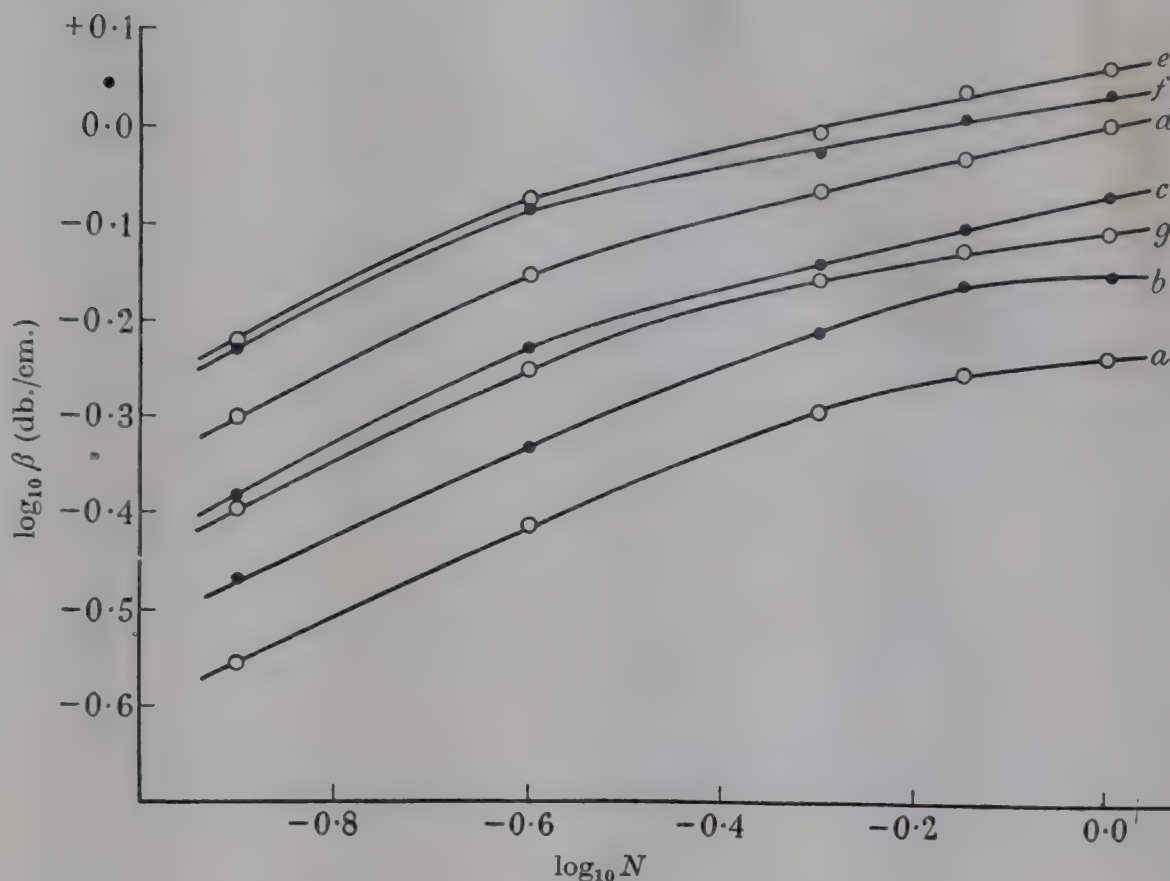


FIGURE 3. Attenuations from potassium bicarbonate.

curve	a	b	c	d	e	f	g
air/gas ratio	2.44	2.67	2.98	3.23	4.00	4.25	4.50
reversal temperature ($^{\circ}$ K)	2015	2060	2100	2140	2190	2170	2120

This is not seriously at variance with the actual value of 4.32 eV for potassium. The two points well off the full line of figure 4 correspond with air-rich flames, which show a somewhat greater curvature than the gas-rich flames of the same temperature in figure 3.

The halogen salts of potassium were then compared and table 2 shows clearly the effects observed. At all concentrations and in all the flames the fluoride and chloride gave results not detectably different from those for bicarbonate, whereas the bromide gave a lower, and the iodide a still lower attenuation. In a general way these may be explained by saying that the hydrogen of the flame gases removes fluorine and chlorine to such an extent that there are insufficient halogen atoms left to pick up a substantial number of electrons, whereas the bromine and iodine are not so much

reduced in concentration. The iodide results are plotted in figure 5 in the same manner as for bicarbonate and do not show the limiting slope of 0.5 so clearly, or the pro-

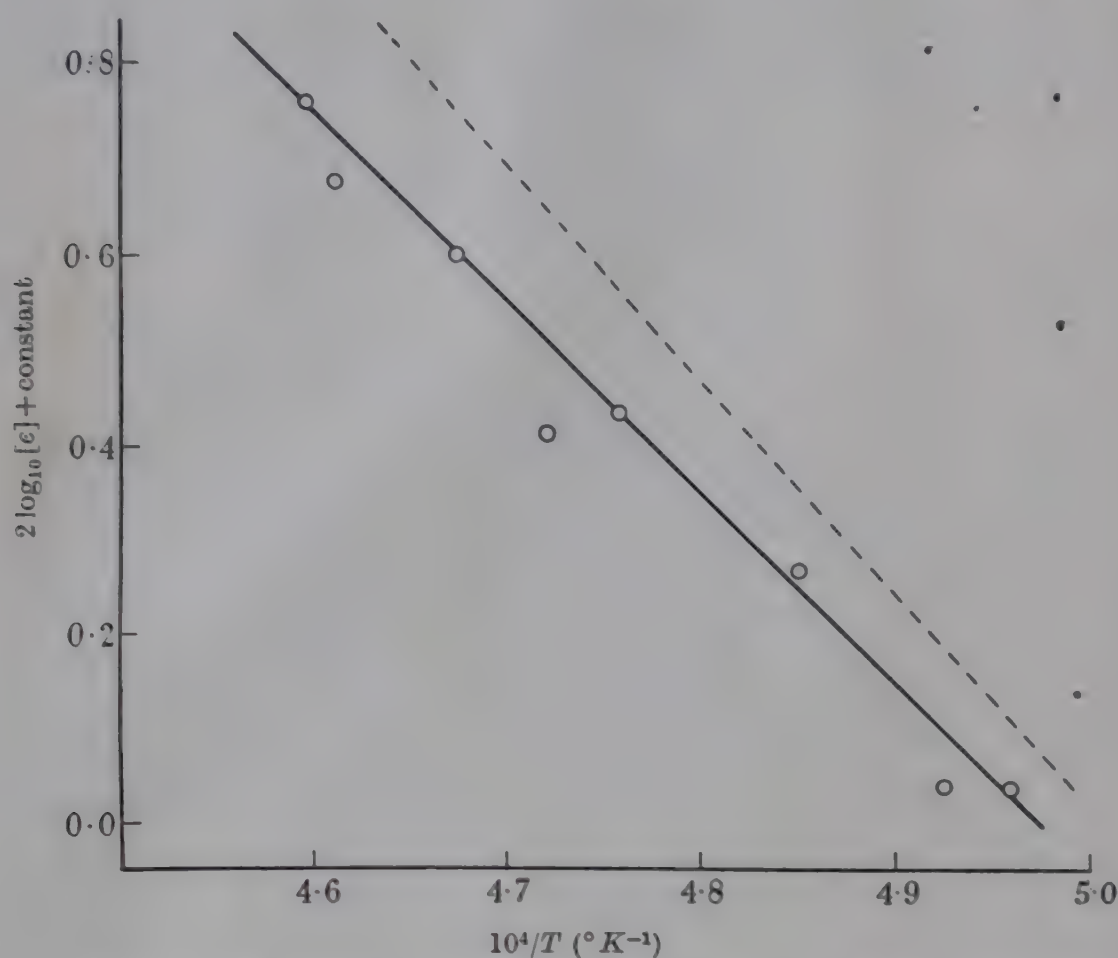


FIGURE 4. Plot for determination of ionization potential of potassium: —, slope = 4.0 ± 0.2 eV; ---, slope = 4.4 ± 0.2 eV.

TABLE 2. ATTENUATIONS OF POTASSIUM HALIDES (IN DB. CM.)

flame composition air/gas	normality of solution	KHCO ₃	KF	KCl	KBr	KI
2.44:1	N	0.59	0.59	0.59	0.54	0.50
2015° K	N/2	0.52	0.52	0.52	0.47	0.46
	N/8	0.29	0.28	0.29	0.27	0.24
3.58:1	N	1.20	1.18	1.20	1.08	0.99
2190° K	N/2	0.99	0.99	0.96	0.90	0.83
	N/8	0.60	0.60	0.60	0.60	0.50
4.50:1	N	0.79	0.79	0.76	0.71	0.63
2120° K	N/2	0.71	0.71	0.71	0.60	0.50
	N/8	0.40	0.40	0.40	0.38	0.35

nounced curvature towards higher concentrations, but are more nearly straight lines of lower slope than the Saha equation predicts.

An attempt to obtain a quantitative explanation of these results will now be made.

(b) Discussion of potassium bicarbonate results

In this and subsequent discussions a number of expressions will be deduced in which certain parameters $K_1 \dots K_7$ will be introduced. If the whole flame system is in chemical equilibrium, then these parameters will have the full significance of equilibrium constants of the reactions to which they refer. If such equilibrium does not hold good, then the parameters will signify the ratios of concentrations which define them, and expressions in which they appear may be used to describe the

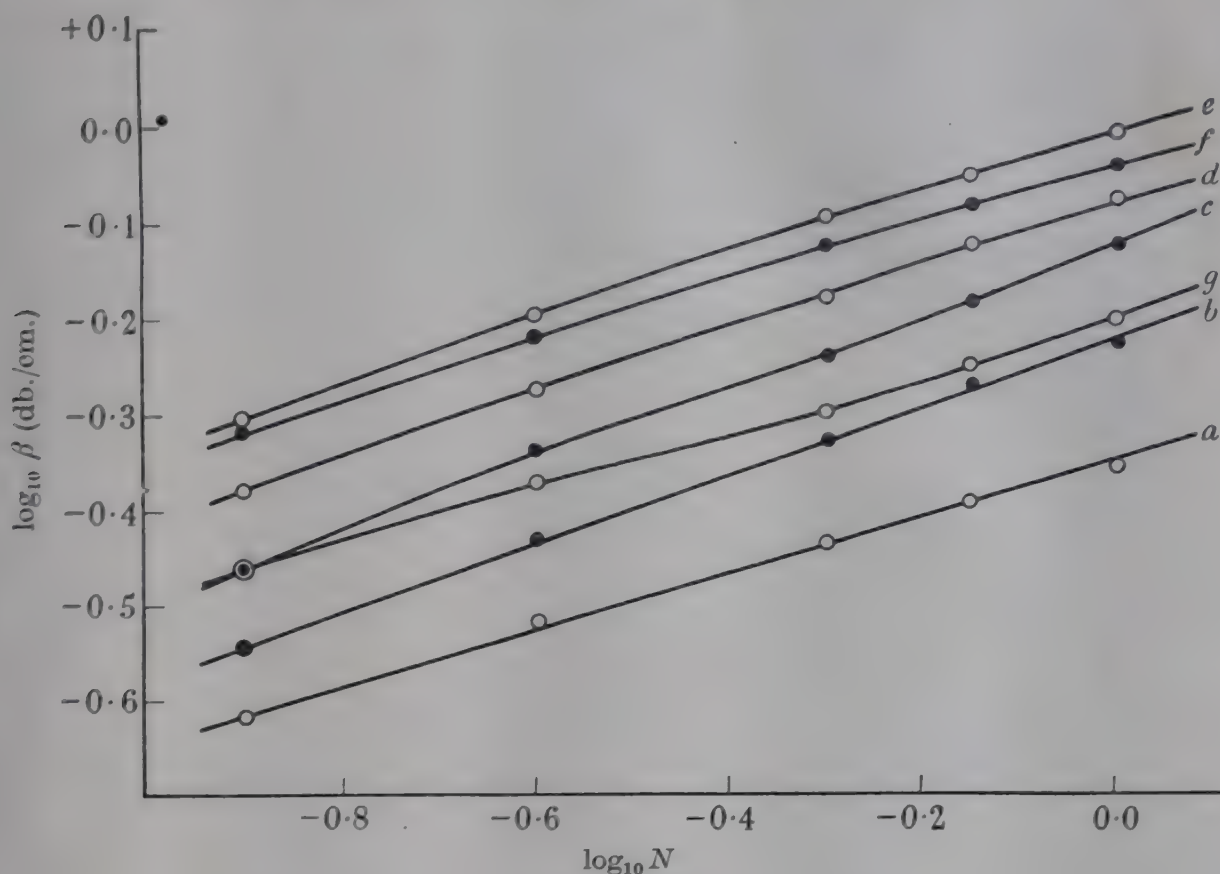


FIGURE 5. Attenuations from potassium iodide.

curve	a	b	c	d	e	f	g
air/gas ratio	2.44	2.67	2.98	3.23	4.00	4.25	4.50
reversal temperature ($^{\circ}$ K)	2015	2060	2100	2140	2190	2170	2120

concentrations of chemical substances involved. It is considered* that two types of disequilibrium should be distinguished, even if the distinction is somewhat arbitrary: isothermal disequilibrium in which a given pair of reactions have not yet reached their equilibrium condition, but in which all the reactants are in temperature equilibrium with their surroundings, and a disequilibrium arising from the fact that one or more reactants are not in temperature equilibrium. In the latter case it is possible that a quasi-chemical equilibrium is reached if both forward and backward reactions are rapid; it would tend to true chemical equilibrium as the temperature equilibrium became established. In this case it is possible to consider the above-mentioned parameters as representing quasi-equilibrium conditions at different temperatures for different reactions in the same system. This method of approach to the complicated chemical system of the flame was used, since any approach involving

velocity constants would be much more difficult and arbitrary. The reversal temperatures will, in the first place, be taken as a norm from which any divergences from equilibrium will be measured.

If a concentration $[X]$ of an alkali metal A is added, and neither A nor A^+ reacts with any constituent of the flame gases, we have the simple result

$$A \rightleftharpoons A^+ + e, \quad K_1 = \frac{[A^+][e]}{[A]}, \quad [A^+] = [e],$$

and since only a small fraction is ionized here then we may say $[A] \gg [A^+]$ and $[A] = [X]$. Then

$$[e]^2 = K_1[X]. \quad (1)$$

Values of K_1 may be obtained with accuracy from theoretical data since the ionization potentials of the alkali metals are accurately known. Values of $\log_{10} K_1$ against $1/T$ are plotted in figure 6, calculated from the theoretical expression

$$K_p = \frac{(2\pi m k T)^{3/2}}{h^3} \left(\frac{kT}{e} \right) e^{-\chi/kT},$$

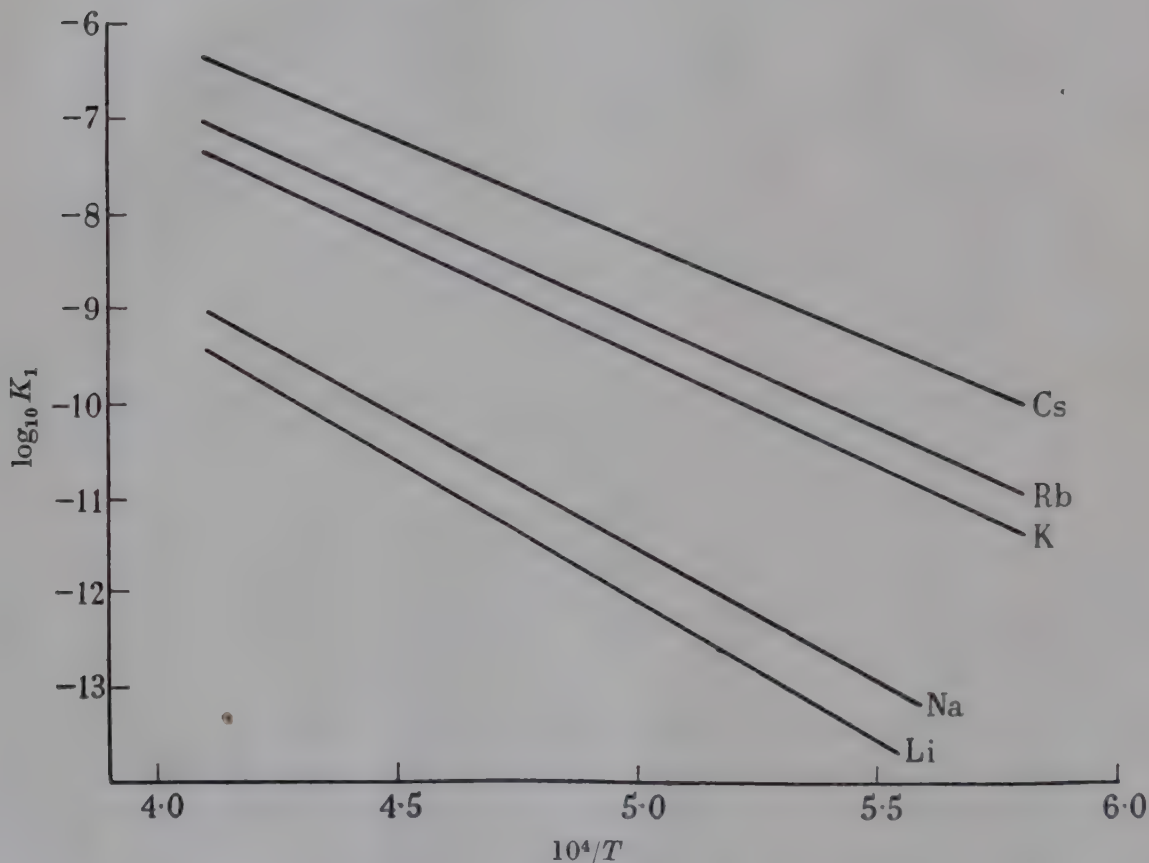
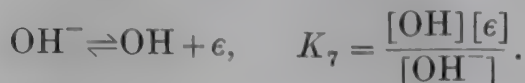
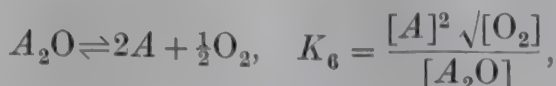


FIGURE 6. Calculated values of K_1 for alkali metals.

where K_p is K_1 in c.g.s. units, m is electron mass and χ the ionization potential of A . The values in figure 6, used in this work, are in atmosphere units, since all concentrations are considered as partial pressures in these units. Equation (1) has been found to hold good for potassium bicarbonate at low concentrations (and also for fluoride and chloride), except that K_1 observed is always about one-fifth of the calculated value.

In order to explain this divergence and the failure of (1) at higher concentrations it is necessary to take into account other equilibria which may affect the value of $[A]$.



To solve this system the conservation equations for mass and electric charge are required:

$$[X] = [A] + [A^+] + [AOH] + 2[A_2O],$$

$$[A^+] = [\epsilon] + [OH^-].$$

Neglecting $[A^+]$ compared with $[A]$ in the first of these, and substituting for all A compounds and OH^- in terms of $[\epsilon]$, a quadratic in $[\epsilon]^2$ results:

$$\frac{\psi}{K_1^2} [\epsilon]^4 (1 + \phi')^2 + [\epsilon]^2 \left(\frac{1 + \phi}{K_1} \right) (1 + \phi') - [X] = 0, \quad (2)$$

where
$$\phi = \frac{[OH]}{K_5}, \quad \phi' = \frac{[OH]}{K_7} \quad \text{and} \quad \psi = \frac{2\sqrt{[O_2]}}{K_6},$$

which solves to
$$[\epsilon]^2 = -\frac{a}{2c} [1 - (1 + 4c[X])^{\frac{1}{2}}], \quad (3)$$

where
$$a = \frac{K_1}{(1 + \phi)(1 + \phi')} \quad \text{and} \quad c = \frac{\psi}{(1 + \phi)^2}.$$

If $4c[X] \ll 1$ this gives as an approximation

$$[\epsilon]^2 = a[X](1 - c[X]), \quad (3a)$$

and if $c[X]$ is negligible
$$[\epsilon]^2 = a[X]. \quad (3b)$$

$[OH]$ and $[O_2]$ may be considered constants, since they are buffered by large quantities of H_2O and H_2 in the flame gases. (3b) agrees with the low concentration results if $(1 + \phi)(1 + \phi') \sim 5$, while at higher concentrations the effect of the parameter c is to cause a curvature as is observed in figure 3. Quantitative agreement between (3) and figure 3 may be obtained by choosing suitable values of c . The values of $c[X]$ required to give this agreement vary from 1.0 to 1.5 for N solutions ($[X] = 3.2 \times 10^{-5}$ atm.) in proceeding from a gas-rich flame at 2015° K to stoichiometric at 2190° K, and become larger (up to 2.0) on the air-rich side. This is to be expected if oxide formation is important, since as $[O_2]$ increases, ψ will increase and also c , unless offset by changes in ϕ . It will be seen that the curve g of figure 3 lies below c of the same diagram, although its reversal temperature is higher, which is understandable if $[OH]$ is larger on the air-rich side for a given temperature, since $(1 + \phi)(1 + \phi')$ will be larger and hence a smaller. Table 1 shows that $[OH]$ does in fact increase in this direction, so that the prediction that equation (3b) will give rise to a lower attenuation on the air-rich side is fulfilled in the range where that equation is applicable.

Figure 4 will give the correct ionization potential only if c and $(1 + \phi)(1 + \phi')$ are independent of temperature. A correction for c may be applied using the figures which fit figure 3 to the N/8 solutions ($[X] = 4 \times 10^{-6}$ atm.) where the effect is small, and gives the dotted line of figure 4, with a slope giving 4.4 eV for the ionization potential.

The question of whether the reactions involving AOH , OH^- and A_2O are expected to be important at these temperatures is difficult to decide, since the thermodynamic data required for evaluation of $K_{5,6,7}$ are scanty and incomplete. In the case of K_5 there is doubt about the partition function of AOH and the energy difference ΔE_0 of the reaction. Bichowsky & Rossini (1936) quote a figure of 32.5 kcal./mole for heat of vaporization, and on the basis of the potassium halides it is reasonable to take 5 kcal./mole for the heat of fusion; hence an approximate value of ΔE_0 may be computed from thermochemical data. K_5 will not vary seriously with the vibration frequencies, shape and size of AOH , and its value has been calculated by introducing three moments of inertia with reduced mass of 10 hydrogen units and radius of gyration 1.5 Å, and of two frequencies of 10^{13} sec. $^{-1}$ and one of 10^{14} sec. $^{-1}$, which are all reasonable values. This gives $K_5 = 8 \times 10^{-3}$ atm. at 2200° K and 1×10^{-3} atm. at 2000° K, resulting in $\phi \equiv [\text{OH}]/K_5 \sim 1$ at all temperatures, using $[\text{OH}]$ at reversal temperatures from table 1. The calculation is very approximate (to an order of magnitude) but is not inconsistent with the experimental ratio of $[\epsilon]^2$ and $K_1[X]$ of about 1:5 which was found.

In the case of K_7 for its value in c.g.s. units at constant pressure we have

$$K_5 = \frac{q_\epsilon q_{\text{OH}}}{q_{\text{OH}^-}} \left(\frac{kT}{e} \right) e^{-\Delta E_0 R/T},$$

where the q is the appropriate partition function. A fair assumption to make is that q_{OH} is not very different from q_{OH^-} , and so the only doubtful quantity is ΔE_0 , the electron affinity of OH. Bichowsky & Rossini (1936) quote -58 kcal./mole for heat of formation of OH^- which would imply an electron affinity of -64 kcal./mole for OH; general considerations with regard to the properties of OH indicate that some error of sign has been made, and it is noteworthy that the same authors have misquoted the sign in one of their references to OH^- (that of Lederle 1932). The best value which seems to be available is one calculated from lattice energies by Gonbeau & Klemm (1935) of 2.1 eV, but even this is rather uncertain. If 2.0 eV is used then K_7 at 2000 and 2200° K respectively is 5×10^{-3} and 1.5×10^{-2} atm., while at 3.0 eV, values of 3×10^{-6} and 8×10^{-5} result. In the former case, using $[\text{OH}]$ from table 1 at reversal temperatures, $\phi' \equiv [\text{OH}]/K_7$ is equal to 0.5 or less, while in the latter $\phi' \sim 100$. All that can be said is that supposition of equilibrium between electrons and hydroxyl is reasonable if the electron affinity is about 2.0 eV but not if it is as high as 3.0 eV, since this would require values of a in (3b) of about 0.01 of K_1 .

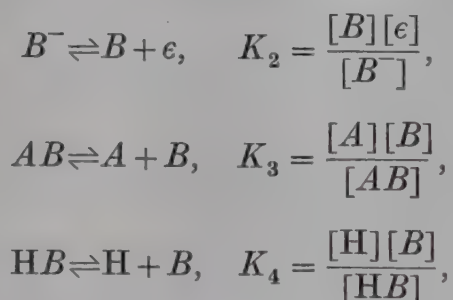
The oxide equilibrium is the most difficult from the data aspect. Nearly all data are lacking on the oxide K_2O , but a rough calculation has been carried out to find whether the results give a reasonable value for the heat of formation of this oxide. If it is taken that $c \equiv 2 \sqrt{[\text{O}_2]}/K_6$ (i.e. neglecting ϕ in $c = \psi/(1 + \phi)^2$), then at 2200° K, K_6 must be about 10^{-7} atm. $^{\frac{1}{2}}$ to explain the observations. Assuming K_2O to be linear

with K-O distances of 2 Å, and with two vibration frequencies of $10^{13} \text{ sec.}^{-1}$ and two of $10^{13.3} \text{ sec.}^{-1}$, then the heat of formation of gaseous K_2O must be about 50 kcal./mole. The heat of formation of the solid from oxygen gas and potassium vapour is computed to be 126 kcal./mole (data of Bichowsky & Rossini 1936), requiring therefore a heat of sublimation of about 75 kcal./mole. This is not unreasonable since, for example, the heats of sublimation of the alkaline earth oxides are of the order of 90 to 110 kcal./mole and those of the potassium halides 40 to 50 kcal./mole. The equilibrium constant varies rapidly only with ΔE_0 and T .

It therefore appears that the ionization produced by adding potassium bicarbonate to a flame is explicable in terms of complete equilibrium between electrons and alkali metal if reasonable account is taken of the possibilities of formation of hydroxide, oxide and hydroxyl ions in equilibrium amounts also at the flame (reversal) temperature. It is obviously impossible to make a more definite statement than this in the absence of complete thermodynamic data, of direct measurements of the concentrations listed in table 1 and of the concentrations of the other substances involved in the above treatment.

(c) Discussion of results for potassium halides

In the case of a halide, a number of additional equilibria may become important:



where B is the halogen, and if different amounts of alkali metal A and halogen B are added the balance equations become

$$\begin{aligned} [X] &= [A] + [A^+] + [AOH] + [AB] + 2[A_2O], \\ [Y] &= [B] + [B^-] + [HB] + [AB], \\ [A^+] &= [e] + [OH^-] + [B^-], \end{aligned}$$

where $[X]$ is total A added and $[Y]$ total B , expressed as partial pressures. It is not possible readily to obtain a general solution of these equations, together with those used previously, unless it is assumed that ionization of A is relatively small, so that $[A] \gg [A^+]$; since the results differ only slightly from those for bicarbonate, where this assumption appeared to be quite valid, it will be taken to be so here, and $[A^+]$ and $[B^-]$ ignored in the equations for $[X]$ and $[Y]$. Further, it will be noted from table 2 that the results for fluoride and chloride are sensibly the same as those for bicarbonate, and so it would appear unnecessary to consider the equilibrium involving $[AB]$ (undissociated halide) in these cases, and hence also the bromide and iodide, which are more readily dissociated. Making these changes the mass balance

equations for $[X]$ and $[Y]$ may be expressed as follows in terms of $[A]$ and $[B]$ and the parameters previously used:

$$[X] = [A](1 + \phi + \psi[A]),$$

$$[Y] = [B](1 + \theta),$$

including also $\theta \equiv [H]/K_4$. From the charge balance expression and the equation for K_1 we may say that

$$[A] = \frac{[\epsilon]^2}{K_1} \left(1 + \phi' + \frac{[B]}{K_2} \right),$$

and substituting in the $[X]$ equation and rearranging we obtain a quadratic in

$$\frac{[\epsilon]^2}{K_1} \left(1 + \phi' + \frac{[B]}{K_2} \right).$$

$$\frac{[\epsilon]^4}{K_1^2} \left(1 + \phi' + \frac{[B]}{K_2} \right)^2 \psi + \frac{[\epsilon]^2}{K_1} \left(1 + \phi' + \frac{[B]}{K_2} \right) (1 + \phi) - [X] = 0,$$

which solves to

$$\frac{[\epsilon]^2}{K_1} \left(1 + \phi' + \frac{[B]}{K_2} \right) = - \frac{(1 + \phi) \pm \sqrt{\{(1 + \phi)^2 + 4\psi[X]\}}}{2\psi}.$$

Substituting $[B] = [Y]/(1 + \theta)$ and $c = \frac{\psi}{(1 + \phi)^2}$,

$$\frac{[\epsilon]^2}{K_1} \left(1 + \phi' + \frac{[Y]}{K_2(1 + \theta)} \right) = \left(\frac{1 + \phi}{2\psi} \right) [\sqrt{\{1 + 4c[X]\}} - 1],$$

therefore

$$[\epsilon]^2 = \frac{b}{b + [Y]} \frac{a}{2c} [(1 + 4c[X])^{\frac{1}{2}} - 1] \quad (4)$$

and hence a final expression for $[\epsilon]^2$ is obtained, where $b = K_2(1 + \theta)(1 + \phi')$ and $a = K_1/(1 + \phi)(1 + \phi')$, which is the same as equation (3) except that there is now a factor $b/(b + [Y])$. Hence the attenuation given by a halide should diverge more from that given by the bicarbonate as the concentration of added salt increases (for salt alone $[X] = [Y]$). That this is not generally true may be seen from figures 3 and 5, where, for any given pair of curves corresponding to the same flame, although the iodide one has a lower gradient for dilute solutions, it becomes steeper than the bicarbonate one towards higher concentrations. A simple quantitative test of the order of magnitude of the parameter b may be made as follows. Since ϕ' is considered to be small the quantity $K_2(1 + \theta)$ gives an approximate value of b and may be calculated from the values of K_2 and K_4 plotted in figure 7, and obtained from the data of Herzberg (1939) and statistical mechanical equations, together with the reversal temperature values of $[H]$ from table 1.

b calculated for three flames in this manner is given in table 3.

TABLE 3. CALCULATED EQUILIBRIUM VALUES FOR $10^6 K_2(1 + \theta)$

T ($^{\circ}\text{K}$)	F	Cl	Br	I
2015	430	4.2	1.2	1.3
2100	500	5.0	1.2	2.4
2190	300	6.0	2.3	5.3

Two discrepancies with theory are at once apparent. If a halide is to give results identical with those for bicarbonate, then in equation (4) $b \gg [Y]$. This is obviously so with potassium fluoride ($[X]$ for N solutions = 32×10^{-6}) but is not so for the chloride, although experiment gave the same electron population. Secondly, the values for the iodide and bromide are in the wrong order, since the iodide gave the lower attenuation, and once again the values of $K_2(1 + \theta)$ are too low, predicting a much greater difference between the halide and bicarbonate than was actually found. Using the attenuations at $N/8$ and equations (3) and (4), experimental values of b have been found from the results and are given in the short table 4.

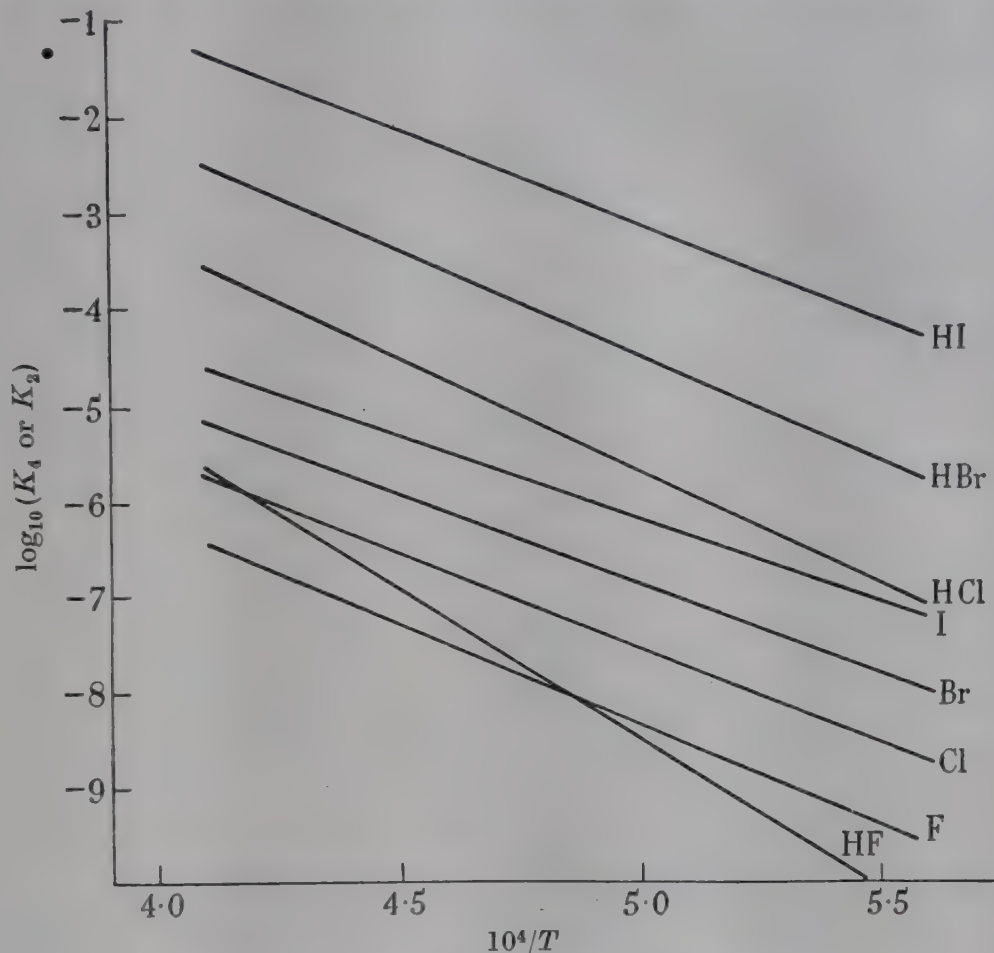


FIGURE 7. Calculated K_2 for halogens, and K_4 for halogen acids.

TABLE 4

T ($^{\circ}$ K)	$10^6 \times b$
	from KI results
2015	24
2100	20
2190	20

These are 10 to 20 times as large as the calculated values of $K_2(1 + \theta)$, which would still not be large enough even if a reasonable allowance was made for $(1 + \phi')$ in the theoretical expression for b . The constancy of these figures, taken with the constancy of $K_2(1 + \theta)$ for each of F, Cl and Br in table 3, and the noticeable increase for I in that table as temperature and air content increase suggest that the error lies in θ

rather than in K_2 —i.e. in K_4 or $[H]$ or both. The idea that this might be concerned with disequilibrium is confirmed when it is realized that a fixed value of b is insufficient to explain the relative effects of iodide and bicarbonate as depicted in figures 5 and 3, which require that b should increase with the concentration of added salt.

It is useful here to consider what the effect of an error in temperature would be while retaining the idea of the maximum of thermal and chemical equilibrium. The study of potassium bicarbonate does not eliminate this possibility owing to absence of data on ϕ , ϕ' and ψ . The relative effect for halides with respect to bicarbonate is depicted in figure 8, which is a plot of $b/(b+[X])$ against temperature, using the $[H]$ concentration at 2190° K reversal temperature (5×10^{-4} atm.) in estimating $b = K_2(1+\theta)$, and $[X]$ for NaX solutions (4×10^{-6} atm.).

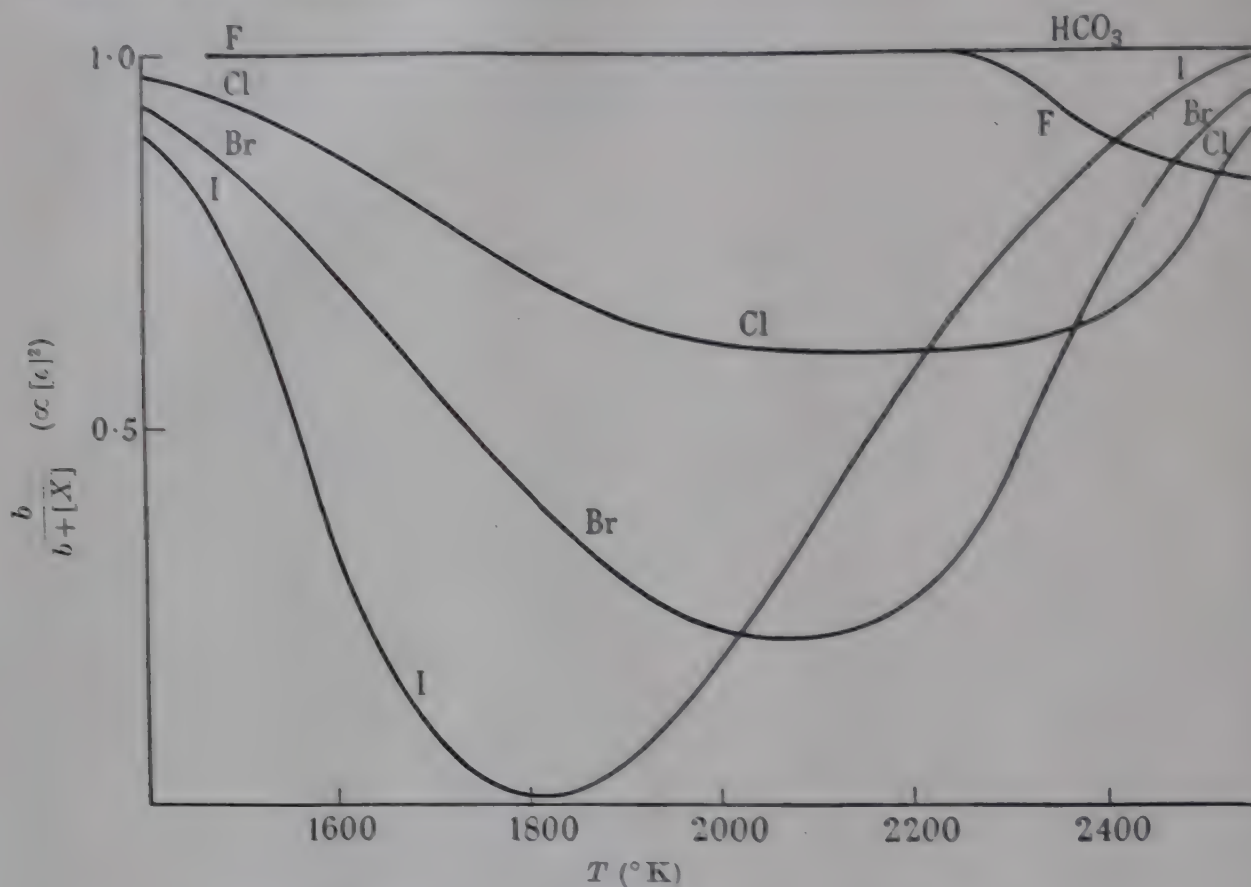
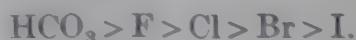


FIGURE 8. Plot of $b/(b+[X])$ against T (°K). $b = K_2(1+\theta)$, $[X] = 4 \times 10^{-6}$ atm., $\theta = [H]/K_4$, $[H] = 5 \times 10^{-4}$ atm.

Each halogen shows a pronounced hollow, these being staggered along the temperature axis. At high temperatures the negative ions are unstable, giving attenuations



whereas at low temperatures the formation of the acids prevents negative-ion formation with attenuations in the order



Under these conditions the order and relative magnitudes of the observed results could only be obtained if K_2 and K_4 corresponded with a temperature some 600 to

800° C below the reversal value. By allowing $[H]$ to increase, approximately the correct behaviour may be obtained at 2190° K (figure 9), but only at a very high $[H]$ level.

Explanation may be sought in various directions, first, the hydrogen atom concentration may be very large—much larger than its equilibrium value; secondly, the effective temperature of both the K_2 and K_4 equilibria might be low, or in the case of K_2 high with K_4 unchanged; finally, there may be sufficient disequilibrium of the isothermal kind in the K_2 or K_4 reactions. The measured value near constancy

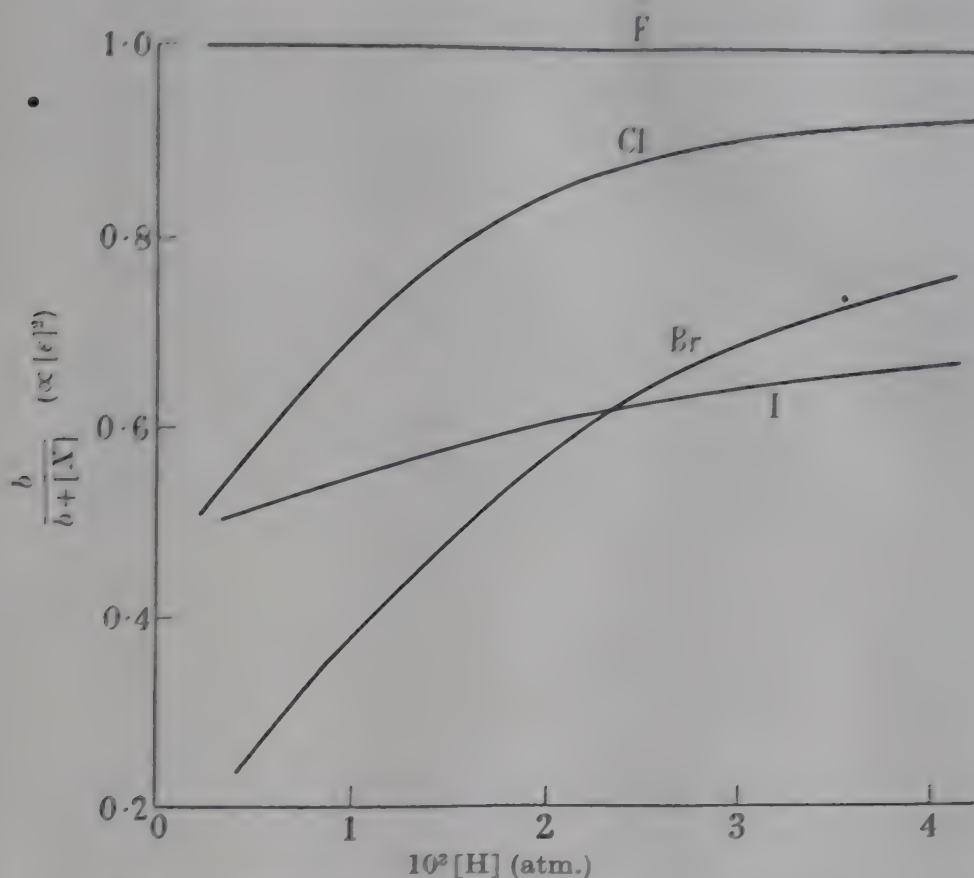


FIGURE 9. Plot of $b/(b + [X])$ against T (°K). $b = K_2(1 + \theta)$, $[X] = 4 \times 10^{-6}$ atm. (N/8 solution), $\theta = [H]/K_4$, 2190° K.

and the calculated variation in θ for the iodide suggest that the error lies in θ rather than in K_2 , as does also the fact that if the K_1 reaction equilibrates rapidly then it is difficult to see why K_2 should not. There is no particular reason for supposing isothermal disequilibrium in the K_4 reaction since there are so many rapid and available chemical processes by which it may equilibrate. This leaves the suggestion of quasi-equilibrium of K_4 , with one or more reactants in a different state of activation than the surroundings, which may also be linked with an abnormally high value of $[H]$.

It is not unreasonable to suppose that a non-equilibrium distribution of energy obtains among the hydrogen atoms, which are well known to be very important chain centres in the oxidations of hydrogen and carbon monoxide taking place in all parts of the flame; at the same time, they are important centres in the hydrogen halogen reactions. If they are on the whole activated—i.e. effectively at a higher temperature

than their surroundings, while the B atoms and HB molecules are at or near temperature equilibrium—then it is to be expected that the rate of formation of HB will be increased while its dissociation will remain at the reversal temperature rate, so that K_4 will be lowered and θ increased. An increase in θ ($= [H]/K_4 = [HB]/[B]$) of 10 to 20-fold is sufficient to explain the discrepancies.

It is not possible to separate out the two effects of concentration and activation of H in the observed effects, and in any case they will be interconnected. In general it appears, however, that complete, or nearly complete, equilibration can occur except in the reactions involving chain centres for the flame gases. It has been supposed that the attenuation method might have been used as a basis for flame-temperature measurement, one of the projected methods involving the relative attenuations given by the various halides. This now appears to be unlikely as a suitable method, on account of the associated anomalies, which would result in meaningless temperatures, as is shown by figure 8.

5. THE EFFECT OF EXTRA HALOGEN

Some very remarkable effects have been observed when extra halogen is added in the potassium salt studies, which have proved to be quite inexplicable on an equilibrium theory. They are shown in figures 10*a* and *b*. Figure 10 has abscissae proportional to $[Y]$, the total iodine or chlorine respectively in *a* and *b*. From 10*a*,

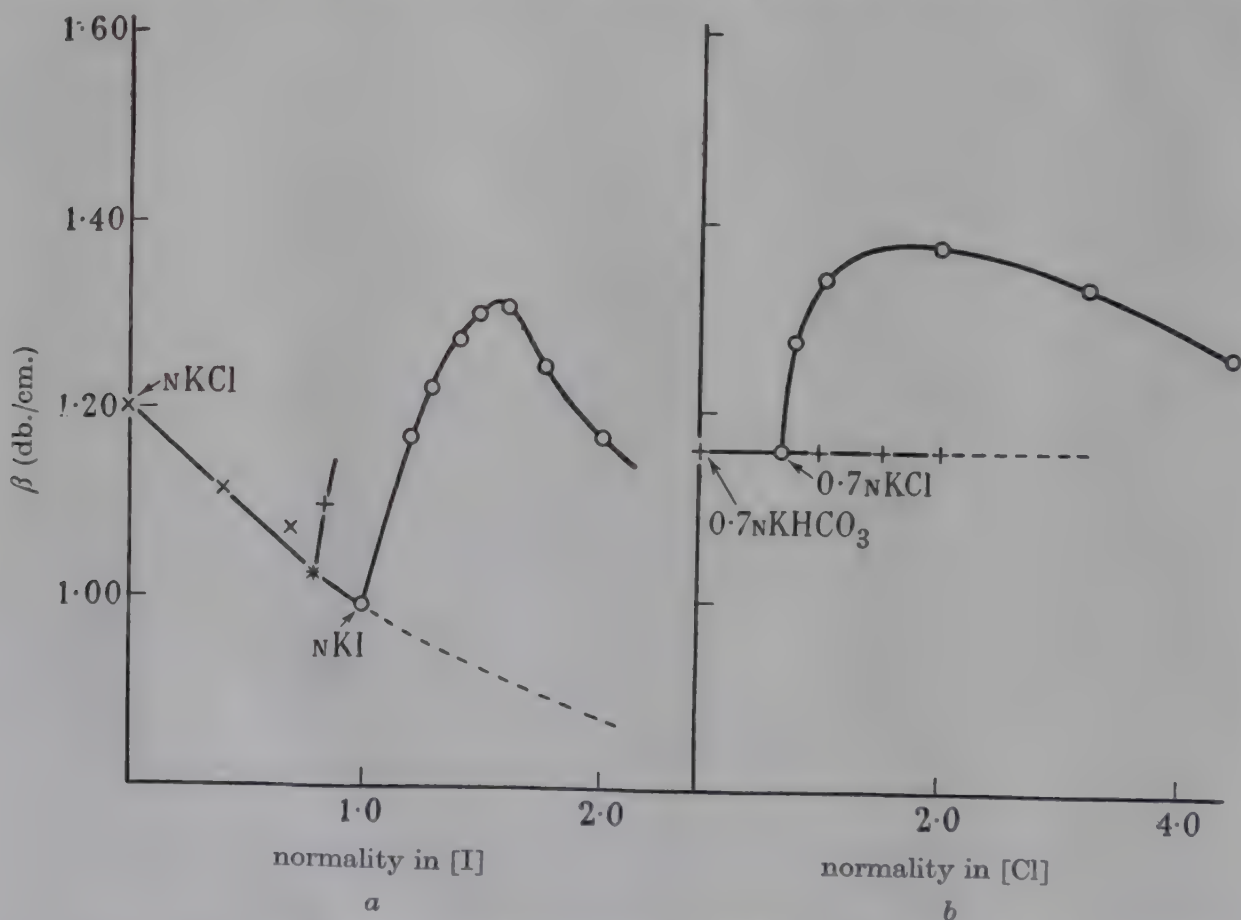


FIGURE 10. Effect of extra halogen. *a*, addition of I_2 to $N-KI$, of $N-KI$ to $N-KCl$, and of I_2 to mixture of $N-KI$ and $N-KCl$. *b*, addition of HCl to $0.7N-KCl$, and of chloral hydrate to $0.7N-KCl$.

mixtures of potassium chloride and iodide show a fairly regular decrease with iodine concentration, as would be expected. Further increase in iodine by using solutions of iodine in potassium iodide shows a sudden discontinuity with a rise of attenuation to a maximum, followed by a sharp fall. That this is a real discontinuity is shown by adding iodine at a point along the chloride/iodide mixture where a similar effect is observed. The theoretical equation (4) is obviously incapable of dealing with these phenomena.

The effect is even more striking in the chloride case of figure 10*b*. Potassium bicarbonate and chloride and mixtures of them all show the same attenuation. Addition of extra chlorine as chloral hydrate gave no detectable change, as would be expected, since for chlorine b of equation (4) is presumably very large. Addition of extra chlorine as hydrochloric acid again gave a marked maximum, and a similar effect was observed if chlorine gas was used, and also in the case of potassium bromide with bromine. This establishes once again the existence of disequilibrium, since the same amount of halogen added in different ways produces different attenuation results. With hydrochloric acid the results are even more interesting in that the whole of the maximum effect lies above the attenuation given by potassium bicarbonate—i.e. either a higher ‘temperature’ in the K_1 equilibrium is established or ϕ , ϕ' , or ψ are modified, probably ϕ or ϕ' , since the effect was also observed in dilute solutions. It was present in all flame-gas compositions to about the same extent, the plotted results being for the stoichiometric flame at 2190° K. Any process causing these results is more likely to be seated in the lower blue zone of flame reaction, since it appears to involve the nature of the halogen compound directly added, and presumably complete equilibration is not attained by the time the gases reach the upper zone where the attenuation is measured. In this connexion it is interesting that all the anomalies are produced by substances which will be available for reaction at a very early stage since they will be present in the unburnt gas stream either as vapours or readily vaporizable substances, whereas the salts will not. It is, nevertheless, felt that the observed effects are not to be ascribed to incomplete vaporization and dissociation from the original spray, since the attenuations did not vary when the wave-guides were moved vertically over a range of about 1 cm. about the mean measuring position, and suitable precautions were taken to remove large droplets.

It is noteworthy that once again an anomaly may be associated with an important chain centre in the main flame reaction—in this case the hydroxyl radical, whose concentration determines ϕ and ϕ' . In this connexion the presence in flames of ClO, BrO and IO, which have been noted by Coleman & Gaydon (1948), may be relevant on account of processes such as



but although these spectra were present in our flames, they had the same intensity in whatever way a given amount of halogen was added. The equilibrium amounts of such substances appear to be negligible in the main reaction scheme.

Although a full explanation of these results is not yet possible without direct measurements of some of the concentrations of other substances in the flames, it is

important to notice that the attenuation method is a powerful one for indicating features of flame reactions which have so far not been indicated by other methods.

6. A COMPARISON OF THE ALKALI CHLORIDES

Finally, a set of readings have been taken with the various alkali chlorides. It was only possible to obtain readings over a wide range of concentration with caesium, rubidium and potassium, since the attenuations with dilute solutions of sodium and lithium salts were too low to be measured accurately. Typical results with a gas-rich, an approximately stoichiometric and an air-rich flame are shown in figure 11. It will be observed that rubidium and caesium do not give the slope of 0.5 corresponding with $[\epsilon]^2 \propto [X]$ in the region of observation; caesium, in particular, in air-rich flames

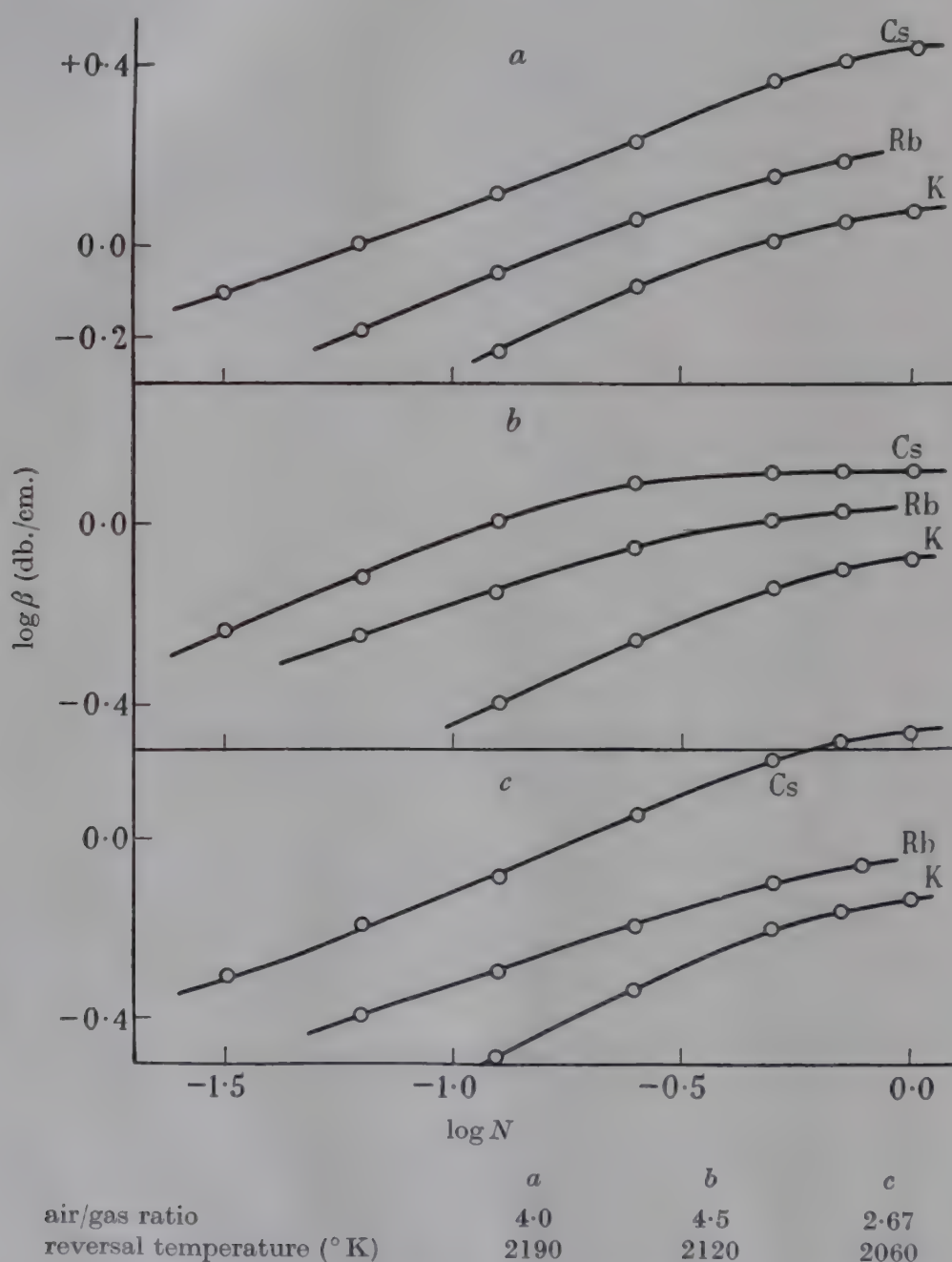


FIGURE 11. Attenuations of potassium, rubidium and caesium chlorides.

shows a tendency to reach a kind of attenuation saturation. A series of results for the various alkali metals are shown in table 5 for a gas-rich flame.

Potassium chloride has been shown to give results which will fit an equation:

$$[\epsilon]^2 = -\frac{a}{2c} [1 - (1 + 4c[X])^{\frac{1}{2}}],$$

where $a = K_1/(1 + \phi)(1 + \phi')$ and $c = \psi/(1 + \phi)^2$, and shows no evidence of undissociated salt or of negative-ion formation. There is no reason to suppose that the other chlorides will show much departure from this, since the dissociation constants are not very different. It is obvious from figure 11 that the rubidium and caesium

TABLE 5. ATTENUATIONS FOR ALKALI CHLORIDES

air/gas 2.67; $T_{\text{rev.}} 2060^\circ \text{K}$

salt	attenuation β db./cm. (N sol.)
LiCl	0.06
NaCl	0.09
KCl	0.71
RbCl	0.90
CsCl	1.70

salts must have higher values of c than potassium on account of the lower slopes of the curves, even for dilute solutions. c may be evaluated from these curves and the equation (3) corrected to the attenuation which would occur in the absence of oxide effect ($c = 0$). ϕ' should be the same for all the metals, and if ϕ is small compared with unity, then a plot of this corrected attenuation ($\log \beta$) against ionization potential should give a slope of $2525/T$, according to the Saha equation expressed as

$$\log_{10} \beta = \frac{-2525}{T} + \text{constant}.$$

This may be illustrated by the results at 2190°K for a stoichiometric flame; values of c of 0.5×10^5 , 1.0×10^5 and 1.5×10^5 give the correct form of the $\log \beta / \log N$ curves up to $0.7N$ for K, Rb and Cs respectively. They do not give the rapid flattening observed with Cs at high concentrations. These c values have been used to correct the observed attenuations for the oxide effect and the logarithm of this corrected β for $N/8$ solutions plotted against V (ionization potential) in figure 12. Points $\circ$ show the corrected attenuations, corresponding to a temperature of $\sim 2500^\circ \text{K}$, while the uncorrected points $\times$ give rise to a temperature of $\sim 3000^\circ \text{K}$. The corrected figure is still 300°C above the reversal temperature, and rather higher than the calculated one, but it still contains the ϕ term of a , which may cause a change in slope if ϕ is not $\ll 1$. If the correct slope (for 2190°K) were to be obtained from these results, it would be necessary for the heats of formation of the hydroxides in the gaseous state to be in the order $\text{CsOH} > \text{RbOH} > \text{KOH}$ with differences of about 1 kcal./mole at $\phi = 1$, but no information on this point is available. Similar results have been obtained for other flames.

An attempt has been made to bring the sodium and lithium results into line by correcting the Cs, Rb and K results for the oxide effect, and assuming it to be small

for Na and Li. Results are given in figure 12 as points + for ∞ solutions, and do not give a good straight line, even allowing for a considerable error in the smaller attenuations. The best straight line, however, agrees more closely with the reversal temperature than do the other lines of figure 12.

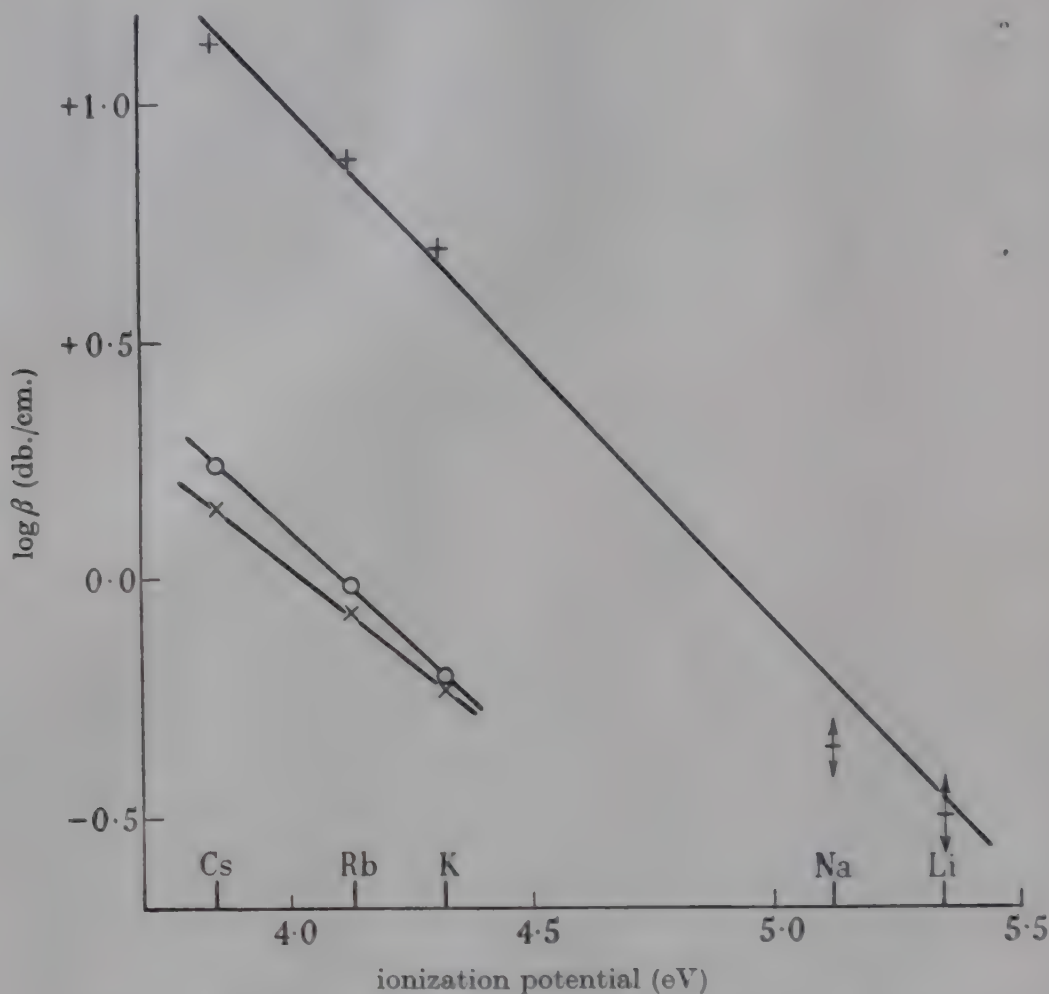


FIGURE 12. Plots of log attenuation against ionization potential.
Slopes; $\times \equiv 3000^\circ \text{K}$, $\text{O} \equiv 2500^\circ \text{K}$, $+\equiv 2300^\circ \text{K}$.

There is therefore no evidence that marked disequilibrium at reversal temperatures occurs for the alkali-ion reaction $A \rightleftharpoons A^+ + e$. It is unfortunate that the parameters ϕ , ϕ' and ψ are not all available from the data, but only a and c , which prevents complete elucidation of the results, and there is no doubt that the present scheme is not capable of dealing with the effects arising at high concentrations of caesium, especially in air-rich flames, but it is considered that the scheme provides a reasonable working hypothesis for handling deviations from the Saha law. The difficulties obviously lie in the complexity of the system, and it is hoped that future work will be possible with simpler flame systems.

7. GENERAL CONCLUSIONS

To sum up, it has been shown that application of the attenuation method to the measurement of electron concentration in flames leads to results which reveal that over a wide range of concentration it is insufficient to consider the simple ionization

reaction of the alkali metal alone, as exemplified by the Saha equation. For a salt not producing any electron acceptors in the flames, the behaviour is reasonably consistent with thermodynamic equilibrium at the measured flame temperature if the possibilities of side reactions with production of alkali metal hydroxide and oxide, and also of OH^- ions are allowed, as far as may be concluded from the limited thermochemical data available.

When salts producing electron acceptors are added (halides) the general trends are consistent with what would be expected from the opposing effects of production of negative halogen ions and of halogen acid. A quantitative explanation in terms of thermodynamic equilibrium proved to be impossible here, however, and the form of the results tends to suggest that this may be brought about by absence of equilibrium or abnormal concentrations of hydrogen atoms in the reactions involving halogen acids. The general effect is to give the relative effects of the four halogens in an order not expected from theory, and to make the chlorides considerably less effective in reduction of electron concentration than would be expected.

These indications of abnormal behaviour with halides are confirmed by addition of extra free or combined halogen, which in some cases produced an increased attenuation, which is in no event consistent with temperature equilibrium of the substances involved. A comparison of the various chlorides without extra halogen, where the halogen appears to have very little effect, allows of a reasonable assessment of the flame temperature as compared with the value measured by the line-reversal method.

The authors wish to express their gratitude to Professor R. G. W. Norrish, F.R.S., for constant advice and encouragement in the course of this research, and to the Royal Society for a grant of apparatus.

REFERENCES

- Belcher, H. & Sugden, T. M. 1950 *Proc. Roy. Soc. A*, **201**, 480.
Bennett, J. A. J. 1927 *Phil. Mag.* **3** (vii), 127.
Bichowsky, F. R. & Rossini, F. D. 1936 *Thermochemistry of chemical substances*. New York: Reinhold.
Coleman, E. H. & Gaydon, A. G. 1947 *Disc. Faraday Soc.* No. 2, 166.
David, W. T. 1935 *Nature*, **135**, 470.
Goubeau, J. & Klemm, W. 1935 *Z. phys. Chem. B*, **36**, 362.
Griffiths, E. & Awberry, J. H. 1929 *Proc. Roy. Soc. A*, **123**, 401.
Herzberg, G. 1939 *Molecular spectra and molecular structure. I, Diatomic molecules*. New York: Prentice-Hall, Inc.
Jones, G. W., Lewis, B., Friauf, J. B. & Perrott, G. St J. 1931 *J. Amer. Chem. Soc.* **53**, 896.
Lederle, E. 1932 *Z. phys. Chem. B*, **17**, 362.
Lewis, B. & von Elbe, G. 1938 *Combustion, flames and explosions of gases*. Cambridge University Press.
Lewis, B. & von Elbe, G. 1943 *J. Chem. Phys.* **11**, 75.
Margenau, H. 1945 *Phys. Rev.* **69**, 508.
Saha, M. N. 1920 *Phil. Mag.* **40**, 472.
Wilson, H. A. 1944 *Modern physics*, p. 321. London: Blackie.

Transient source, doublet and vortex solutions of the linearized equations of supersonic flow

BY W. J. STRANG

*Aeronautical Research Laboratories, Department of Supply
and Development, Fishermen's Bend, Melbourne*

(Communicated by G. Temple, F.R.S.—Received 30 November 1949)

By introducing the Heaviside step function as the time variation of the fundamental source solution of the wave equation of compressible gas flow, a family of 'transient' source and doublet solutions is derived, together with a generalization of the vortex, and their more interesting properties are explored. These results are applicable to unsteady supersonic aerofoil theory and similar problems.

1. INTRODUCTION

The operational calculus was devised by Heaviside as a means of investigating the conditions in an electrical system in the period immediately following the imposition of a disturbance. Such initial motions are classed as 'transient phenomena' and have attracted much attention in recent years, especially in the fields of electricity, servo-mechanisms and aerofoil theory. Heaviside's operational calculus, placed on a rigorous basis by Bromwich and others, has proved a valuable tool for these studies. An auxiliary concept introduced by Heaviside in this connexion is the unit-step function, defined by

$$H(t) = 0 \quad \text{if } t < 0, \\ = 1 \quad \text{if } t \geq 0,$$

and the response of a system to a disturbance having this mathematical expression is called its 'indicial admittance', a term derived from electrical theory. If $a(t)$ is the indicial admittance of a linear system, the response to a disturbance $f(t)$ can be calculated from the integral formula

$$x(t) = f(0)a(t) + \int_0^t \frac{df(\tau)}{d\tau} a(t-\tau) d\tau,$$

the integral on the right being known by a variety of names of which 'superposition integral' is the most explanatory. These ideas have received close attention in the course of the development of the operational calculus, and the necessary mathematical background exists.

Difficulties sometimes arise from the discontinuity in $H(t)$ at $t = 0$, and a general method of discussing these is to define the step function as the limit of a set of analytic functions, e.g.

$$H(t) = \frac{1}{2} + \frac{1}{\pi} \lim_{\alpha \rightarrow \infty} \tan^{-1} \alpha t.$$

In technical application these questions are usually passed over, with satisfactory results. The 'derivative' of $H(t)$ with respect to t , variously known as the unit-

impulse function or Dirac's δ -function, is the function $\delta(t)$ which is infinite at $t = 0$ and zero everywhere else, in such a way that

$$\int_{-\infty}^{+\infty} \delta(t) dt = 1.$$

Formal differentiation of the above expression for $H(t)$ gives

$$\delta(t) = \frac{1}{\pi} \lim_{\alpha \rightarrow \infty} \frac{\alpha}{1 + \alpha^2 t^2},$$

and it is sometimes convenient to take

$$\delta(t) = \lim_{\alpha \rightarrow \infty} \begin{cases} \alpha & \text{for } 0 \leq t \leq 1/\alpha, \\ 0 & \text{for all other } t. \end{cases}$$

A thin aerofoil in an *incompressible* fluid constitutes a linear system, and the indicial admittance functions have been determined for three unit-step disturbances: change of incidence without rotation, change of angular velocity, entry into a sharp-edged gust.

A thin aerofoil in a *compressible* fluid constitutes a non-linear system, for which the superposition integral is invalid, but a linear approximation can be obtained by taking the speed of sound as a constant a , and this is known to yield satisfactory results in a great many cases. To this approximation, indicial admittance functions for two-dimensional supersonic aerofoils have been obtained, and a recent paper by Evvard (1948) contains a result of wide generality, having implicit in it admittance functions for three-dimensional supersonic wings. Admittance functions for a family of supersonic delta wings have been obtained by the author and will be published shortly.

In all these cases the mathematical problem is: a linear partial differential equation for the velocity potential is given, with the boundary conditions of steady flow at infinity, and specified derivatives of the velocity potential normal to a given plane over a portion of it. The time variation of the given derivatives involves unit-step functions. It is required to find the velocity potential, or the pressure distribution which depends on its derivatives. In a previous paper (Strang 1948) the step function was introduced into the *source solutions*. This made possible a treatment which keeps the mathematical work linked to the physical problem without excessive strain on the imagination. The idea invites development, and this paper describes a number of such source and doublet solutions, which are of interest in themselves. Applications will be given elsewhere.

The well-known steady source solution of the wave equation has an infinity over the Mach cone, and the steady doublet solution has an infinity of higher order. This is a major difficulty in the general theory, and occasioned Hadamard's device of the 'finite part' of an infinite integral. Robinson in England (Robinson 1947) and Heaslet & Lomax in America (Heaslet & Lomax 1947) have approached the problems of supersonic aerofoil theory in this way with considerable success, but the application of this and related ideas to the sources and doublets of this paper (which have other discontinuities arising from the step function) has yet to be examined.

For the present, it is legitimate to regard the use of these results as a way of deriving likely solutions, whose validity may be established *a posteriori*. It is worth remembering that the difficulties in the general theory can often be evaded in specific cases by considering space distributions of sources and doublets, for which the infinities are of lower order. Some line distributions are particularly convenient.

2. THE WAVE EQUATIONS OF LINEARIZED COMPRESSIBLE FLOW

We begin by considering the propagation of small disturbances in a medium of density ρ_0 , at pressure p_0 , initially at rest. If the motion is irrotational, the velocity is derivable from a velocity potential defined by $u = -\text{grad } \phi$, and it is well known that ϕ satisfies the equation

$$\nabla^2 \phi = \frac{1}{a^2} \frac{\partial^2 \phi}{\partial t^2}, \quad (2.1)$$

where a is the speed of sound, a constant in the linearized theory. To this approximation, the pressure change is

$$\Delta p = p - p_0 = \rho_0 \frac{\partial \phi}{\partial t}. \quad (2.2)$$

In a frame of reference moving in the negative x -direction with speed U , equation (2.1) transforms to

$$a^2 \nabla^2 \phi - U^2 \frac{\partial^2 \phi}{\partial x^2} = \frac{\partial^2 \phi}{\partial t^2} + 2U \frac{\partial^2 \phi}{\partial x \partial t}, \quad (2.3)$$

which is the linearized equation of compressible flow. This paper is concerned only with the cases for which $U > a$, that is to say, with cases of supersonic flow, although subsonic flow can be treated similarly.

The formula for the pressure change when the potential is referred to the moving co-ordinate system is

$$\frac{\Delta p}{\rho} = \frac{\partial \phi}{\partial t} + U \frac{\partial \phi}{\partial x}. \quad (2.4)$$

Equations (2.1) and (2.3) are described as wave equations, which implies that they have 'real' characteristics in the (x, y, z, t) space. These characteristics are the loci $\theta = \text{const.}$, where

$$\left(\frac{\partial \theta}{\partial x}\right)^2 + \left(\frac{\partial \theta}{\partial y}\right)^2 + \left(\frac{\partial \theta}{\partial z}\right)^2 = \frac{1}{a^2} \left(\frac{\partial \theta}{\partial t}\right)^2,$$

or

$$(a^2 - U^2) \left(\frac{\partial \theta}{\partial x}\right)^2 + a^2 \left(\frac{\partial \theta}{\partial y}\right)^2 + a^2 \left(\frac{\partial \theta}{\partial z}\right)^2 = \left(\frac{\partial \theta}{\partial t}\right)^2 + 2U \left(\frac{\partial \theta}{\partial x}\right) \left(\frac{\partial \theta}{\partial t}\right),$$

respectively, and each characteristic is a possible branch locus of an integral surface. They may be thought of as wave fronts in (x, y, z, t) space. In the source and doublet solutions which follow, there are surfaces dividing the fields of flow into regions in which the potentials have different analytic forms. It may be verified that these surfaces are sections of the characteristics by the planes $t = \text{const.}$

3. SOURCE SOLUTIONS

3.1. Stationary sources

Equation (2.1) has the two fundamental solutions

$$\phi = \frac{1}{r_s} f(at - r_s) \quad \text{and} \quad \phi = \frac{1}{r_s} f(at + r_s),$$

where r_s is the distance from a fixed point, and all possible solutions can be derived from these. We are particularly concerned with two special cases of the first of these:

(i) the potential

$$\phi = \frac{Va}{4\pi r_s} \delta(at - r_s), \quad (3.1.1)$$

which represents a source which emits volume V of fluid as a pulse at $t = 0$, and

(ii) the potential

$$\phi = \frac{V}{4\pi r_s} H(at - r_s), \quad (3.1.2)$$

which represents a source which emits fluid at a constant rate V per unit time for $t > 0$.

An account of these solutions has already been given (Strang 1948), and the discussion will not be repeated here. It is to be noted that for $r_s < at$, the potential of (3.1.2) is a solution of Laplace's equation also, so this source is an 'incompressible' source as far as the sphere of radius at .

Questions of aerofoil theory are concerned with sheets of sources, and line sources are useful in such applications. Line-source solutions have been derived (Strang 1948) for the case of an infinite line, but we shall extend the results to the case of a terminated line. The influence of the terminus is limited to the surface and interior of a sphere of radius at , and it is sufficient to consider the case of a line source extending from 0 to $+\infty$ along the x -axis. We suppose that the source emits volume V per unit length per unit time for $t > 0$, and the velocity potential is readily obtained by integrating (3.1.2). The solution has cylindrical symmetry, so we introduce the cylindrical co-ordinate $r_c^2 = y^2 + z^2$. Denoting the potential of (3.1.2) by $\phi_1(r_c, x)$ at any fixed time $t > 0$, the new potential is

$$\phi = \int_0^\infty \phi_1(r_c, \xi - x) d\xi = \int_{-x}^\infty \phi_1(r_c, \xi) d\xi.$$

The integrand is non-zero only if $|\xi| < T$, where $T = +\sqrt{(a^2 t^2 - r_c^2)}$, and this makes the effective limits of integration $\int_{-x}^{+T}$ in region A of figure 1 and $\int_{-T}^{+T}$ in region B .

The resulting potential is

$$\left. \begin{aligned} \text{(i)} \quad \phi &= \frac{V}{4\pi} \left\{ \cosh^{-1} \left(\frac{at}{r_c} \right) + \sinh^{-1} \left(\frac{x}{r_c} \right) \right\} \quad \text{if } x^2 + r_c^2 < a^2 t^2, \\ \text{(ii)} \quad \phi &= \frac{V}{2\pi} \cosh^{-1} \left(\frac{at}{r_c} \right) H(at - r_c) \quad \text{if } x^2 + r_c^2 > a^2 t^2, x > 0, \\ \text{(iii)} \quad \phi &= 0 \quad \text{outside these limits.} \end{aligned} \right\} \quad (3.1.3)$$

The potential (ii) is the two-dimensional solution and represents radial flow to a cylindrical wave front. Equation (i) indicates spreading of the flow at the cut end as shown in figure 1. The potential for a line source terminated finitely at both ends is derivable from this by superposition. For any finite length of line source, the end spheres will intersect after a sufficiently long time, and still later each sphere will contain the whole line.

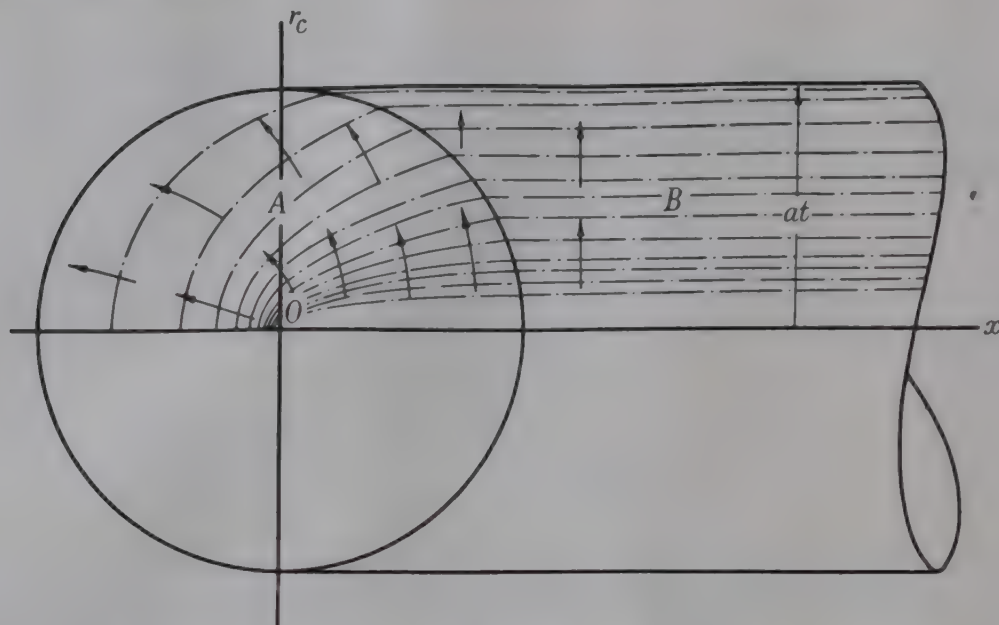


FIGURE 1. Stationary line source from $x = 0$ to $x = +\infty$.
 — — —, equipotential; $\rightarrow$, direction of flow.

The equations of this section are all solutions of equation (2.1). They transform directly into solutions of equation (2.3), and then represent sources drifting downstream with the fluid. Since the potential depends only on x , r_c and t , the solution surfaces can be drawn in 3-space, and this affords an opportunity of illustrating the role of the characteristics as surfaces of demarcation. For the case of cylindrical symmetry, equation (2.1) becomes

$$\frac{\partial^2 \phi}{\partial r_c^2} + \frac{1}{r_c} \frac{\partial \phi}{\partial r_c} + \frac{\partial^2 \phi}{\partial x^2} = \frac{1}{a^2} \frac{\partial^2 \phi}{\partial t^2},$$

and the equation of the characteristics is

$$\left(\frac{\partial \theta}{\partial r_c} \right)^2 + \left(\frac{\partial \theta}{\partial x} \right)^2 - \frac{1}{a^2} \left(\frac{\partial \theta}{\partial t} \right)^2 = 0.$$

Among the solutions of this are the cone $x^2 + r_c^2 - a^2 t^2 = \theta = 0$ and the pair of tangent planes $a^2 t^2 - r_c^2 = \theta = 0$. The diagram (figure 2) shows these, and it will be seen that any plane $t = \text{constant} > 0$ cuts these in a configuration similar to figure 1.

3.2. Supersonic source solutions

Supersonic point source. The solution of equation (2.1) given by equation (3.1.1) can be made to yield moving source solutions, since a moving source which is emitting fluid continuously can be represented by a chain of pulses. In particular, we can calculate the velocity potential of a source, moving through the fluid with constant

velocity $U > a$, which emits a constant volume V per unit time for $t \geq 0$ and does not operate for $t < 0$. It is convenient to use a co-ordinate system moving with the source, and the velocity potential is then a step-function solution of equation (2.3) having the same fundamental nature as (3.1.2).

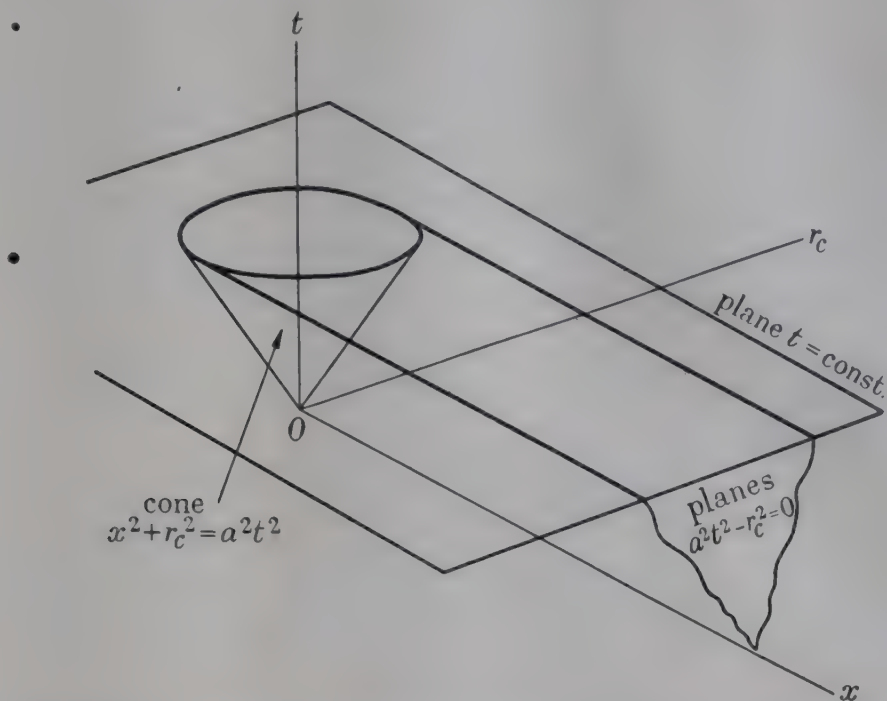


FIGURE 2. Stationary line source from $x = 0$ to $x = \infty$ showing the characteristic surfaces $a^2 t^2 - r_c^2 = 0$ and $x^2 + r_c^2 = a^2 t^2$.

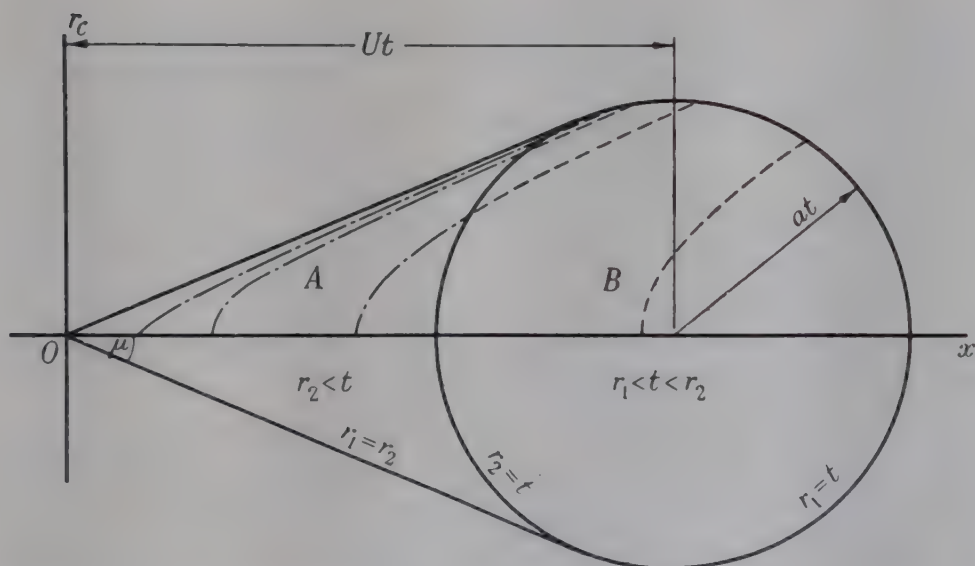


FIGURE 3. Supersonic point source. —, equipotential in A ; ---, equipotential in B .

After time t , the disturbance fills the plummet-shaped region shown in figure 3, and it is convenient to introduce the following generalizations of certain quantities that were found useful in the paper already mentioned (Strang 1948). Considering a field point $P(x, y, z)$, define τ_1 and τ_2 as the roots of the equation

$$(U^2 - a^2)\tau^2 - 2Ux\tau + (x^2 + y^2 + z^2) = 0,$$

selecting $\tau_1 \leq \tau_2$ so that

$$\tau_1 = (Ux - aR)/(U^2 - a^2), \quad \tau_2 = (Ux + aR)/(U^2 - a^2),$$

where

$$R = +\frac{1}{a} \sqrt{\{a^2 x^2 - (U^2 - a^2) r_c^2\}}.$$

The quantity

$$T = +\sqrt{(a^2 t^2 - r_c^2)}$$

will also appear frequently. Both T and R have the dimensions of length, and the τ 's the dimensions of time. A pair of co-ordinates (τ_1, τ_2) defines a ring of axially symmetric points P , which all have the same value of ϕ , and, in fact, $a\tau_1$ and $a\tau_2$ are the radii of the two spheres centred on the x -axis which pass through P . The loci $\tau_1 = t$ and $\tau_2 = t$ together make up the sphere in figure 3 and the locus $\tau_1 = \tau_2$ is the Mach cone.

Viewing the moving source as a chain of pulses, we are led to the following integral for the velocity potential,

$$\phi = \frac{Va}{4\pi} \int_0^t \frac{1}{\varpi} \delta(a\tau - \varpi) d\tau,$$

in which $\varpi^2 = (x - Ut)^2 + r_c^2$. The integrand is non-zero only if $\tau = \tau_1$ or τ_2 and the integral depends on the relative magnitudes of τ_1, τ_2 and t . In region A of figure 3 it is*

$$\phi = V/2\pi R, \quad (3.2.1)$$

which is, of course, the solution for a *steady* supersonic point source. In B it is

$$\phi = V/4\pi R. \quad (3.2.2)$$

The equipotential surfaces are hyperboloids of revolution, homothetic with respect to the origin. At the spherical surface, a given equipotential changes from one hyperboloid to another, a reminder that the supersonic point source can have no physical existence.

The result given by equations (3.2.1) and (3.2.2) taken together stands in the same relation to the transient theory as equation (3.2.1) alone does to the steady theory, but it is not convenient to use, and it is not easy to discuss the validity of operations on a potential of this kind. Many of the difficulties disappear when one- and two-dimensional distributions of sources are considered.

Supersonic line source. Integration is a process which smooths out discontinuities, and one would expect line sources to be easier to work with than point sources. A two-dimensional line source of finite length moving at right angles to its length has been examined. It appears that such a source is not a very useful concept, because its velocity potential involves a large number of domains in which different analytic expressions hold, and because awkward discontinuities persist near the ends. On the other hand, a line source moving 'point-foremost' has proved very

* It is necessary to use the result that

$$\int_{-\infty}^{+\infty} g(x) \delta[f(x)] dx = \sum_i g(x_i) \left| \frac{1}{f'(x_i)} \right|,$$

summed for all the roots x_i of $f(x) = 0$.

useful indeed, and has a number of features which commend it. Such a source will be called a *needle source* to distinguish it from the other kind, and it is sufficient to deal with the case of a uniform distribution of supersonic point sources along the positive x -axis.

The potential has cylindrical symmetry, and it is fairly easy to see on physical grounds that the field of flow will divide into the regions shown in section in figure 4. The actual expressions are obtained without much difficulty by integrating the results obtained above, and are quoted without proof. In the conical space A of figure 4

$$\phi = \frac{V}{2\pi} \cosh^{-1} (x \tan \mu / r_c). \quad (3.2.3)$$

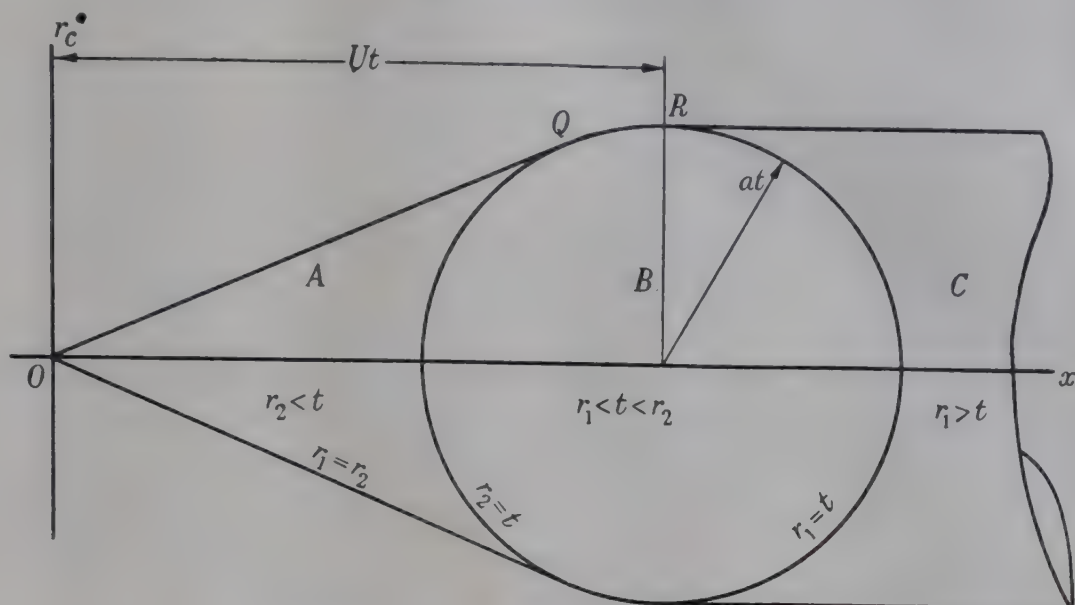


FIGURE 4. Needle source or doublet.

This result applies also to a steady needle source ($t \rightarrow \infty$). The equipotentials are cones through the origin. In the spherical space B ,

$$\phi = \frac{V}{4\pi} \{ \cosh^{-1} (at/r_c) + \cosh^{-1} (a\tau_1/r_c) \}, \quad (3.2.4)$$

and in the cylindrical space C ,

$$\phi = \frac{V}{2\pi} \cosh^{-1} (at/r_c), \quad (3.2.5)$$

which is the same as for a stationary line source—a kind of cylindrical expansion.

The continuity of this potential deserves comment. The discontinuity over the Mach cone, which is a feature of the point source, has disappeared, and the solution is finite everywhere except at the source itself. It has also a finite discontinuity along the zone QR , but it is otherwise continuous. Moreover, only three regions are involved, which is very moderate in this work.

The pressure distribution for this source is important. According to the linearized theory, the pressure change p corresponding to a velocity potential satisfying equation (2.3) is given by

$$\Delta p / \rho_0 = \frac{\partial \phi}{\partial t} + U \frac{\partial \phi}{\partial x}.$$

Applying this, we get

$$\left. \begin{aligned} \Delta p/\rho_0 &= UV/2\pi R \quad \text{in } A, \\ \Delta p/\rho_0 &= \frac{V}{4\pi} \left\{ \frac{a}{T} + \frac{U}{R} \right\} \quad \text{in } B, \\ \Delta p/\rho_0 &= Va/2\pi T \quad \text{in } C. \end{aligned} \right\} \quad (3.2.6)$$

This pressure is infinite but integrable over the surface of the Mach cone and the cylinder, and has finite discontinuities over the surface of the sphere. Nevertheless, it is not intractable, and it is this that makes a needle source so useful. It means that it is possible to calculate the pressure distribution directly, without reference to the velocity potential, and as it happens that this leads to easier integrals there is a double saving.

Gust source. Prominent in the aerodynamics of unsteady motion are problems involving entry into gusts. The needle source is not suited to such problems and its place is taken by another type of line source which I shall call a gust source. This consists of a semi-infinite line moving in the direction of its length, consisting of sources which start emitting as they cross a certain front located in the fluid, and it will be clear that it differs from a needle source by a *stationary* line source from the front to infinity.

As before, we refer the motion to axes fixed to the tip of the source, and from this moving viewpoint the gust source consists of a uniform distribution of supersonic point sources along the positive x -axis, which start emitting consecutively at time x/U . To obtain the potential, write $(x - Ut)$ for x in equation (3.1.3) (because we are referring the motion to a moving co-ordinate system) and subtract from equations (3.2.4) and (3.2.5). The potential of the gust source is found to be

$$\left. \begin{aligned} \phi &= \frac{V}{2\pi} \cosh^{-1}(x \tan \mu/r_c) \quad \text{in region } A \text{ of figure 4,} \\ \phi &= \frac{V}{4\pi} \{ \cosh^{-1}(a\tau_1/r_c) - \sinh^{-1}(\overline{x - Ut}/r_c) \} \quad \text{in } B, \\ \phi &= 0 \quad \text{in } C. \end{aligned} \right\} \quad (3.2.7)$$

The corresponding pressures are

$$\left. \begin{aligned} \Delta p/\rho_0 &= UV/2\pi R \quad \text{in } A, \\ \Delta p/\rho_0 &= UV/4\pi R \quad \text{in } B. \end{aligned} \right\} \quad (3.2.8)$$

It will be seen that in region A the solution is the same as for a needle source but, since $\phi = 0$ in C , the whole disturbance fills the plummet-shaped region drawn in figure 3 for a point source.

4. DOUBLET SOLUTIONS

4.1. Stationary doublets

Equation (3.1.2) gives the potential for a point source having an emission rate which is a step function of t . Let there be such a source of strength V at $(0, 0, \delta)$ and

another of strength $(-V)$ at $(0, 0, -\delta)$, and let $\delta \rightarrow 0$ and $V \rightarrow \infty$ in such a way that $2V\delta \rightarrow M$. This process gives us a new potential

$$\begin{aligned}\phi &= -\frac{Ma}{4\pi} \frac{\partial}{\partial z} \left\{ \frac{1}{r_s} H(at - r_s) \right\} \\ &= \frac{Ma \sin \theta}{4\pi} \frac{1}{r_s^2} \{ H(at - r_s) + r_s \delta(at - r_s) \},\end{aligned}\quad (4.1.1)$$

where $\sin \theta = z/r_s$. This is the potential of a doublet with step-function emission. For $r_s < at$ this is exactly the same as the velocity potential of a doublet in classical hydrodynamics, and this is a consequence of the fact that $\phi = V/4\pi r_s$ satisfies Laplace's equation. The disturbance fills a sphere of radius at , the velocity being infinite at the surface. Figure 5 shows the equipotentials and the direction of flow.

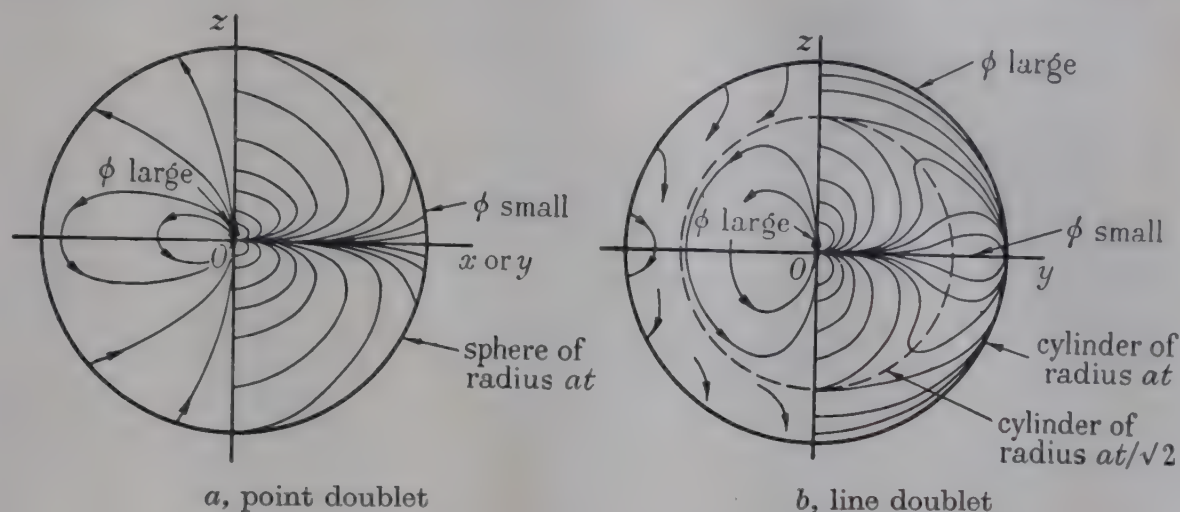


FIGURE 5. Stationary doublets. Directions of flow are shown to the left and equipotentials to the right of each figure.

By differentiating (3.1.3) or by integrating (4.1.1) the following expression is obtained for the potential of a two-dimensional line doublet,

$$\begin{aligned}\phi &= \frac{M \sin \theta}{2\pi r_c} \frac{H(at - r_c)}{\sqrt{1 - (r_c/at)^2}} \\ &= \frac{M}{2\pi} \frac{z}{y^2 + z^2} \frac{at}{T} H(at - r_c).\end{aligned}\quad (4.1.2)$$

The cross-section of this generalized doublet is shown in figure 5. The radial component of velocity for $r_c < at$ is

$$u_r = -\frac{\partial \phi}{\partial r_c} = \frac{M \sin \theta}{2\pi a^2 t^2} \left\{ \frac{2 - (at/r_c)^2}{[1 - (r_c/at)^2]^{\frac{3}{2}}} \right\},$$

which is zero if $r_c/at = 1/\sqrt{2}$, i.e. there is no radial component of velocity across the cylinder of radius $at/\sqrt{2}$. Since the cylinder itself is expanding, this does not mean that there is a barrier in the fluid, but it does mean that the instantaneous flow converges on or diverges from this surface without crossing it. The advancing wave has a 'head and tail' rather like the pulse solution for a stationary line source, i.e. it is a pressure wave, followed by a rarefaction zone, if $z > 0$, and vice versa if $z < 0$.

A line doublet of finite length is the next to be considered, but there is no real loss of generality in supposing it to be semi-infinite, occupying the positive x -axis, with the doublet axis in the z -direction. The velocity potential is most readily derived by differentiating (3.1.3). Using cylindrical co-ordinates in which $r_c^2 = y^2 + z^2$, $\tan \theta_c = z/y$, the doublet potential comes out to be

$$\phi = \frac{M \sin \theta}{4\pi r_c} \left\{ \frac{1}{\sqrt{\{1 - (r_c/at)^2\}}} + \frac{1}{\sqrt{\{1 + (r_c/x)^2\}}} \right\} \quad \text{if } x > 0, \\ \text{or} \quad \phi = \frac{M \sin \theta}{4\pi r_c} \left\{ \frac{1}{\sqrt{\{1 - (r_c/at)^2\}}} - \frac{1}{\sqrt{\{1 + (r_c/x)^2\}}} \right\} \quad \text{if } x < 0, \quad (4.1.3)$$

in the domain for which $x^2 + r_c^2 < a^2 t^2$. If $x^2 + r_c^2 > a^2 t^2$ and $x > 0$ the result of (4.1.2) applies. The change of sign in (4.1.3) can be avoided by using spherical polar co-ordinates, but this is inconvenient in other respects.

4.2. Supersonic doublets

Supersonic point doublet. A point doublet moving with supersonic speed and resulting from the coalescence of two 'pulse' sources differs in no way from a stationary doublet. On the other hand, if the sources are 'step sources' we expect the characteristic features of transient point sources, viz. a sphere and a Mach cone enveloping the disturbance. Thus we look for a doublet solution of equation (2.3), a companion solution to the point source, and it is most easily obtained by differentiating the source solutions, when the following results are obtained.

In the conical space A of figure 3, we get

$$\phi = -\frac{M \cot^2 \mu}{2\pi} \frac{z}{R^3}, \quad (4.2.1)$$

and in the spherical space B

$$\phi = -\frac{M \cot^2 \mu}{4\pi} \frac{z}{R^3}. \quad (4.2.2)$$

The peculiarity of the equipotential surfaces which change their values at the surface of the sphere remains, and in addition we have an infinity of order ∞^3 over the Mach cone, where $R = 0$. It is this infinity which causes the difficulty in the general theory, and which Hadamard's ideas are designed to combat.

Supersonic line doublet. We can conceive of a 'needle doublet', and it is interesting to compare it with Lighthill's doublet (Lighthill 1944). His doublet is not a point doublet but a line of doublets moving point foremost with supersonic speed, and we would expect our needle doublet to reduce to his when $t \rightarrow \infty$. We shall find that this is indeed the case.

The velocity potential is most readily derived by differentiating the source solution. In region A , figure 4, the solution is

$$\phi = \frac{M}{2\pi} \frac{z}{y^2 + z^2} \frac{x}{R}, \quad (4.2.3)$$

which is the same as Lighthill's equation (16). In region B we have, by differentiating equation (3.2.4),

$$\phi = \frac{M}{4\pi} \frac{z}{y^2 + z^2} \left\{ \frac{at}{T} - \frac{x}{R} \right\}, \quad (4.2.4)$$

and in C ,

$$\phi = \frac{M}{2\pi} \frac{z}{y^2 + z^2} \frac{at}{T}. \quad (4.2.5)$$

This solution has an infinity on the Mach cone of the same order as that of the simple point source, and it is comparatively easy to discuss distributions of needle doublets. In applications, the difficulty remains of finding doublet distributions to satisfy stated boundary conditions.

5. GENERALIZED RECTILINEAR VORTEX

In the hydrodynamical theory of thin aerofoils in steady and unsteady motion, vortices have played a prominent part. There are reasons why the concept of a vortex is less convenient in supersonic theory, and very little use was made of vortices until the recent work of Robinson (1947), and Heaslet & Lomax (1947), from which it appears that supersonic vortex theory has its merits.

A familiar result from the theory of the potential flow of an incompressible fluid is that a uniform distribution of doublets over a plane surface is equivalent to a vortex ring around the boundary. As a particular case, a uniform half-plane of doublets, bounded by a straight line and the line at infinity, is equivalent to a rectilinear vortex. The potential of this vortex may be written

$$\phi = \frac{M}{2} \left\{ n + \frac{\theta}{\pi} \right\},$$

where n is an integer, arising from the periodicity, and θ is an angular co-ordinate around the line. The vortex has the following properties:

(i) apart from an irrelevant constant, the vortex potential is invariant for rotation of the axes about the vortex itself, i.e. the plane of the doublets cannot be detected from the vortex potential,

(ii) there is no radial flow,

(iii) the tangential flow velocity is proportional to $1/r$.

The question arises whether a uniform distribution of step-function doublets (equation (4.1.1)) may be equivalent to some generalization of a vortex. To investigate this, we take a uniform distribution of doublets over the half-plane $z = 0, y > 0$, and a convenient starting point is the line doublet solution. Writing the potential of (4.1.2) as $\phi_1(y, z)$ the potential for a half-plane is

$$\begin{aligned} \phi &= \int_0^\infty \phi_1(\eta - y, z) d\eta \\ &= \int_0^\infty \phi_1(\eta, z) d\eta \pm \int_0^y \phi_1(\eta, z) d\eta, \end{aligned}$$

the $\pm$ depending on the sign of y . The first integral reduces to $\frac{1}{2}M(\frac{1}{2}+n)$ and the second to $\frac{1}{2}M\frac{1}{\pi}\tan^{-1}(yat/zT)$, giving the potential

$$\phi = \frac{M}{2} \left\{ \frac{1}{2} + n \pm \frac{1}{\pi} \tan^{-1}(yat/2T) \right\}. \quad (5.1)$$

This is the potential of the generalized vortex. It is multi-valued like the usual vortex potential, and increases by M in a circuit of the origin, this being the strength of the vortex. If we let $r/at \rightarrow 0$,

$$\phi \rightarrow \frac{M}{2} \left\{ n \pm \frac{\theta}{\pi} \right\}.$$

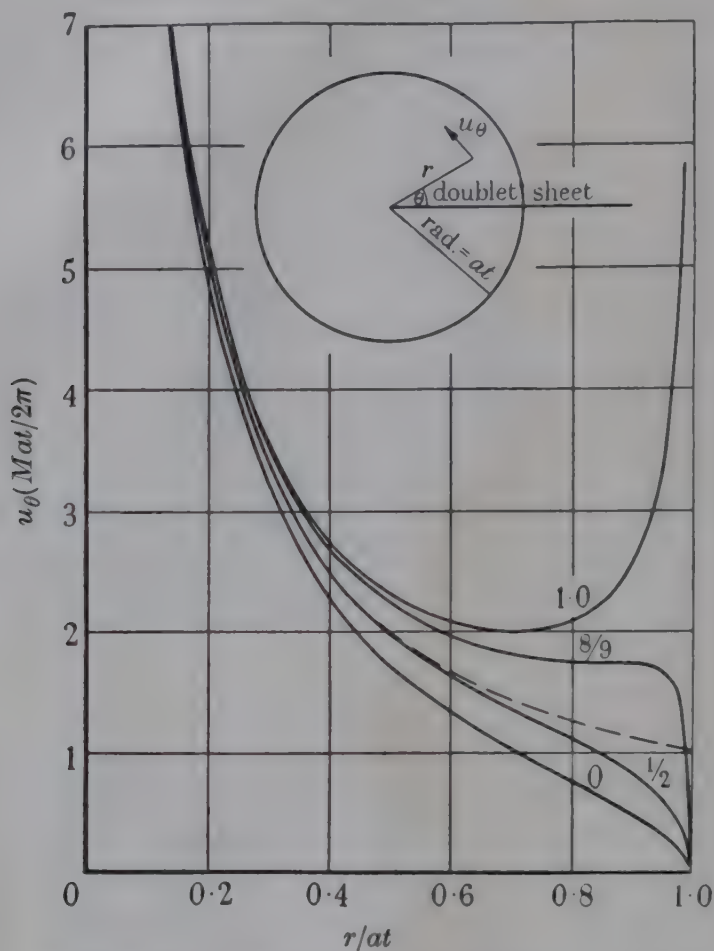


FIGURE 6. Distribution of tangential velocity in a generalized vortex. The numbers on the full curves represent values of $\sin^2 \theta$. ---, incompressible.

So this potential approaches that of the hydrodynamical vortex at large values of the time, or near the core of the vortex at all values of the time. The tangential velocity is

$$u_\theta = -\frac{1}{r_c} \frac{\partial \phi}{\partial \theta} = \frac{Mat}{2\pi} \frac{1}{r_1} \frac{\sqrt{(1-r_1^2)}}{1-r_1^2 \sin^2 \theta},$$

where $r_1 = r_c/at$, and this reduces, of course, to

$$\frac{M}{2\pi r_c} \quad \text{as } t \rightarrow \infty.$$

On the other hand, u_θ differs widely from this value near $r_c = at$, and some graphs of the velocity distribution are shown in figure 6. The orientation of the doublet

sheet from which the vortex is derived remains significant in the solution, and the flux across it is

$$[u_\theta]_{\theta=0} = \frac{Mat}{2\pi} \sqrt{\{(1/r_1^2) - 1\}}.$$

The situation may be described by saying that the vortex is polarized. It is possible to solve problems of the unsteady motion of plane surfaces in a compressible fluid by finding distributions of these generalized vortices for which the flux satisfies given boundary conditions, but the integral equations arising are not easy to solve.

The fact that the doublet sheet remains significant in the step-function vortex potential, but not in the incompressible vortex potential, is somewhat surprising at first sight, but there is an interesting parallel to be drawn with a better-known phenomenon.* Let there be a piston in a long straight tube containing fluid. If motion be imparted to the piston, and the fluid be incompressible, the velocity will be everywhere the same, and the position of the piston cannot be detected from a knowledge of the velocity. If the fluid be compressible, waves will travel through the fluid, originating from the piston, and its position is derivable from the velocity distribution. These features are peculiar to unsteady motion—a steady compressible fluid vortex solution exists which is not polarized.

This work was carried out at the Aeronautical Research Laboratories, Fishermen's Bend, Melbourne, and is published by permission of the Commonwealth Department of Supply and Development, Australia.

REFERENCES

- Evvard, J. C. 1948 *Tech. Notes Nat. Adv. Comm. Aero., Wash.*, no. 1699.
 Heaslet, M. A. & Lomax, H. 1947 *Tech. Notes Nat. Adv. Comm. Aero., Wash.*, no. 1515.
 Lighthill, M. J. 1944 *Rep. Memor. Aero. Res. Comm., Lond.*, no. 2001.
 Robinson, A. 1947 *College of Aeronautics, Cranfield, Rep.* no. 9.
 Strang, W. J. 1948 *Proc. Roy. Soc. A*, **195**, 245.

* This comparison was suggested by Dr J. R. Green.

Transient lift of three-dimensional purely supersonic wings

By W. J. STRANG

*Aeronautical Research Laboratories, Department of Supply
and Development, Melbourne*

(Communicated by G. Temple, F.R.S.—Received 30 November 1949)

A theory for the transient lift of two-dimensional supersonic wings given in a previous paper (Strang 1948) is extended to cover three-dimensional wings of the purely supersonic class. It is found that a simple function of two variables can be defined which can express all the pressure distributions of this and related problems. Charts are given for calculating the unsteady lift of wings with plan-forms made up of a finite number of straight lines, especially delta wings.

For delta wings, the important parameters are Mach number and chord-distance travelled, the sweep-back being less significant. It is concluded that plan-forms of small aspect-ratio are the most advantageous from the point of view of gust-loading, on account of their greater chord-length for given area.

1. INTRODUCTION

The study of the non-stationary aerodynamic forces on an aerofoil moving with constant forward velocity can proceed along alternative lines. In the one, prior consideration is given to the forces arising when an aerofoil performs harmonic oscillations throughout all time. The forces arising in arbitrary motion are then obtainable as a Fourier integral, or its equivalent. In the other, the case of a step-function disturbance (e.g. a sudden change of incidence) is considered first, and the general case is obtained as a superposition integral. Both methods have been used for the case of an incompressible fluid, and complete equivalence of the results has been established.

Motion in a compressible fluid is essentially a non-linear problem, to which the Fourier and superposition integral techniques do not apply, and very poor progress has been made with it as such. To make the problem manageable, recourse is made to linearization of the equations in a manner which makes the speed of sound a constant. The use of the linear approximation is not very satisfactory, for its predictions are not always sustained by experiment in that field of moderate supersonic Mach numbers where it should hold best. A notable example is the theoretical existence of negative torsional damping at Mach numbers between 1 and $\sqrt{2}$, discovered by Possio (1937) and apparently at variance with the facts (Bratt & Chinneck 1947). Some writers have advanced the view that better results are to be expected with a three-dimensional theory, especially if the wings have 'subsonic leading edges', i.e. if the sweep-back exceeds the complement of the Mach angle. Progress is likely to be made by comparing the results of the linearized theory with those of experiment, and working on the discrepancies. A start in this direction has been made by Jones (1947).

The linearized subsonic theory is closely analogous to incompressible theory, owing to the predominance of the Kutta-Joukowski condition, and similar methods

could be used. The same methods have been extended to the supersonic realm (Robinson 1947, 1948; Heaslet & Lomax 1947), but the differences are fundamental, and other methods are more usual. The oscillations of two-dimensional supersonic aerofoils were treated by Possio as far back as 1937, and the first transient solutions for the two-dimensional supersonic wing are those of Schwarz (1943). His paper seems not to be well known, but an English translation is available in the M.A.P. Volkenrode series. Work on the three-dimensional unsteady theory is just beginning. The exposition of Garrick & Rubinow (1947) is applicable to the unsteady motion of wings with supersonic leading and trailing edges. In their examples, however, they did not make this application, except to the case of a rectangular wing, for which it is invalid. The solution for an oscillating delta wing with subsonic leading edges—essentially a more difficult problem—has been obtained by Robinson (1948) in terms of Lamé functions, while Evvard (1948) has indicated a method of dealing with this and more general problems. The details are not yet worked out, and no numerical results are offered in either case.

In a previous paper (Strang 1948) the transient form of the two-dimensional theory was developed on the basis of line-source solutions with step-function time variation of the emission rate. It involved considerations of an elementary kind throughout, and can be extended to three-dimensional cases. This extension is the work which follows. The analysis is applicable to any purely supersonic aerofoil, but the charts presented are designed to permit rapid treatment of those cases in which the aerofoil is made up of straight lines. The general form of the results is not unexpected, but the possibility revealed herein of expressing all the complicated pressure distributions of this and related problems in terms of one simple function discloses a unity that was not anticipated.

2. THE BOUNDARY-VALUE PROBLEM

The linearized wave-equation for the velocity potential of irrotational disturbances in a supersonic stream is

$$(a^2 - U^2) \frac{\partial^2 \Phi}{\partial x^2} + a^2 \frac{\partial^2 \Phi}{\partial y^2} + a^2 \frac{\partial^2 \Phi}{\partial z^2} = \frac{\partial^2 \Phi}{\partial t^2} + 2U \frac{\partial^2 \Phi}{\partial x \partial t}, \quad (2.1)$$

where a is the speed of sound (a constant to this approximation) and U is the speed of the free stream in the direction of x -increasing. It is required to find a velocity potential satisfying this equation, subject to the boundary condition that the flow is to be tangential at the surface of the aerofoil. It can be shown that it is consistent with the assumptions of linearized theory to satisfy the condition over the plane $z = 0$, instead of at the surface of the aerofoil, which is supposed to be everywhere near this plane. For a flat-plate aerofoil which does not pitch

$$-\left[\frac{\partial \Phi}{\partial z} \right]_{z=0} = v(t),$$

where v is the vertical velocity of the aerofoil.

It is always possible to satisfy this condition with a distribution of time-dependent doublets in the plane $z = 0$, but a simpler method is available in the case of so-called 'purely supersonic' aerofoils. These are characterized by the fact that the top and bottom surfaces are completely independent, and a source distribution is sufficient. It has been shown by Puckett (1946) and is easily reasoned on physical grounds, that a uniform distribution of sources over the surface of the plate gives a uniform normal velocity component over its surface also. Puckett's proof was for sources invariant with time, but it is readily extended to time-dependent sources. We examine, then, the case of a uniform sheet of sources with strength varying like $v(t)$.

Sudden change of incidence

If we take $v(t) = H(t)$, where $H(t)$ is the unit step-function, i.e.

$$H(t) = 0 \quad \text{if} \quad t < 0,$$

$$H(t) = 1 \quad \text{if} \quad t \geq 0,$$

the potential corresponding to the source distribution is the potential of a flat-plate aerofoil with unit-step change of incidence, in the region $z > 0$. Since the problem is antisymmetric, the negative of the solution applies in the region $z < 0$.

Sharp-edged gust

For a flat-plate aerofoil in steady motion entering a sharp-edged gust, the boundary condition over the aerofoil is, of course,

$$\left[\frac{\partial \Phi}{\partial z} \right]_{z=0} = 0.$$

The free stream has velocity components U in the x -direction and $H(Ut - z)$ in the z -direction. The moving plane $x = Ut$ is the gust front and is a vortex sheet. It is well known that the following set of conditions is equivalent:

- (i) the free stream has only the x -component of velocity, and
- (ii) the boundary condition over the aerofoil is that

$$\left[\frac{\partial \Phi}{\partial z} \right]_{z=0} = H(Ut - x).$$

A velocity potential exists and can be found by postulating a uniform sheet of sources filling the portion of the aerofoil traversed by the gust front.

The velocity potential is convenient for the discussion of the equivalent source distributions, and its nature is of some interest in itself. It is, however, possible to calculate the pressure distribution, which is really an acceleration potential, from the source distribution without further reference to the velocity potential. For many of the applications of aerofoil theory, this is convenient.

3. THE SOURCE SOLUTIONS

A systematic account of the source and doublet solutions of equation (2.1) with step-function emission has been given elsewhere (Strang 1950), and only a summary of the results we require will be given here.

Supersonic point source

The potential

$$\Phi = V/2\pi R, \quad (3.1)$$

where

$$R^2 = x^2 - (y^2 + z^2) \cot^2 \mu,$$

$$\sin \mu = a/U,$$

is well known as the potential of a supersonic point source of strength V . It may be verified that the potential of the supersonic point source with strength $VH(t)$ is

$$\Phi_A = V/2\pi R, \quad \Phi_B = V/4\pi R, \quad (3.2)$$

where the suffixes refer to the corresponding regions of figure 1. This whole potential with $V = 1$ will be designated $\Phi_S(x, y, z, t)$ in subsequent work. The potential is discontinuous across all boundaries shown in the figure.

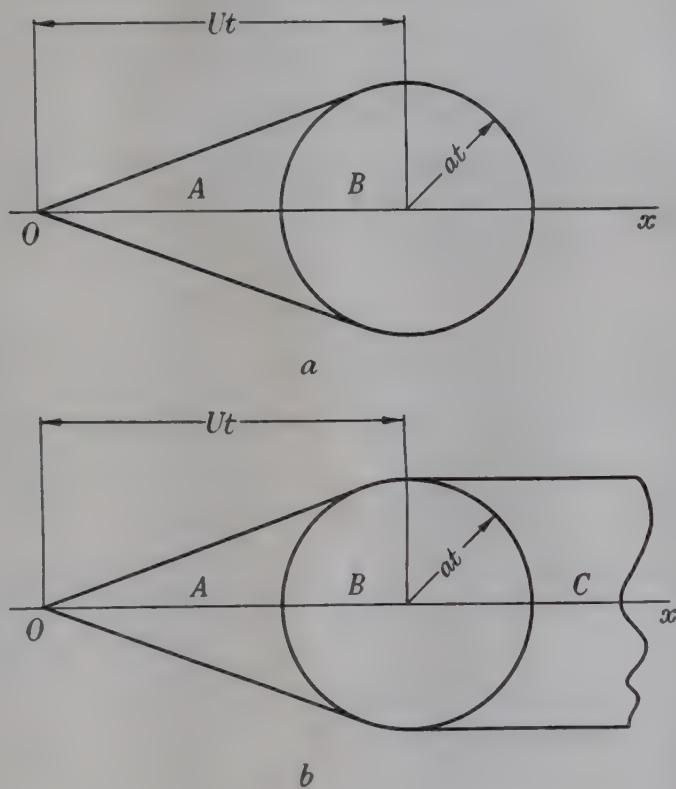


FIGURE 1. Time-dependent supersonic sources. *a*, point or gust source; *b*, needle source.

The potential (3.2) stands in the same relation to the transient theory as equation (3.1) does to the steady theory, but it is not convenient to use, and it is not easy to discuss the validity of operations on a potential of this kind. Many of the difficulties disappear when one- and two-dimensional distributions of sources are considered.

Needle source

A line source moving point-foremost will be called a needle source to distinguish it from a line source across the direction of motion, and it is sufficient to deal with

a uniform distribution of supersonic point sources along the positive x -axis, of strength $VH(t)$ per unit length. The potential is

$$\left. \begin{aligned} \Phi_A &= \frac{V}{2\pi} \cosh^{-1} \left\{ \frac{x \tan \mu}{\sqrt{(y^2 + z^2)}} \right\}, \\ \Phi_B &= \frac{V}{4\pi} \left[\cosh^{-1} \left\{ \frac{at}{\sqrt{(y^2 + z^2)}} \right\} + \cosh^{-1} \left\{ \frac{ar_1}{\sqrt{(y^2 + z^2)}} \right\} \right], \\ \Phi_C &= \frac{V}{2\pi} \cosh^{-1} \left\{ \frac{at}{\sqrt{(y^2 + z^2)}} \right\}, \end{aligned} \right\} \quad (3.3)$$

where $r_1 = (Ux - aR)/(U^2 - a^2)$, and the suffixes refer to figure 1. The potential Φ_A is for the steady flow, and Φ_C represents the radial expansion.

The pressure distribution is given by the following formulae for the pressure increment:

$$\left. \begin{aligned} p_A &= \rho UV/2\pi R, \\ p_B &= \frac{\rho V}{4\pi} \left\{ \frac{a}{T} + \frac{U}{R} \right\}, \\ p_C &= \rho Va/2\pi T, \end{aligned} \right\} \quad (3.4)$$

where $T = +\sqrt{a^2t^2 - (y^2 + z^2)}$. In that which follows, the whole potential (3.3) with $V = 1$ will be written $\Phi_N(x, y, z, t)$, and the pressure (3.4) with $V = 1$ will be written $p_N(x, y, z, t)$. This pressure is infinite but integrable over the surface of the Mach cone and the cylinder, and has finite discontinuities over the surface of the sphere.

Gust source

A uniform distribution of supersonic point sources along the positive x -axis, which start emitting consecutively at time x/U , will be called a gust source. The potential is

$$\left. \begin{aligned} \Phi_A &= \frac{V}{2\pi} \cosh^{-1} \left\{ \frac{x \tan \mu}{\sqrt{(y^2 + z^2)}} \right\}, \\ \Phi_B &= \frac{V}{4\pi} \left[\cosh^{-1} \frac{ar_1}{\sqrt{(y^2 + z^2)}} - \sinh^{-1} \left\{ \frac{x - Ut}{\sqrt{(y^2 + z^2)}} \right\} \right], \end{aligned} \right\} \quad (3.5)$$

where the suffixes A and B refer to figure 1.

The corresponding pressure increment is

$$\left. \begin{aligned} p_A &= \rho UV/2\pi R, \\ p_B &= \rho UV/4\pi R. \end{aligned} \right\} \quad (3.6)$$

In that which follows, the whole potential (3.5) with $V = 1$ will be written $\Phi_G(x, y, z, t)$ and the pressure (3.6) with $V = 1$ will be written $p_G(x, y, z, t)$.

4. SUDDEN CHANGE OF INCIDENCE

Let $x = f_1(y)$ and $x = f_2(y)$ be the leading and trailing edges of the aerofoil under consideration. Since it is purely supersonic $\left| \frac{df_1}{dy} \right|$ and $\left| \frac{df_2}{dy} \right|$ will be everywhere less than $\cot \mu$, where μ is the Mach angle. The boundary conditions require a uniform

distribution of sources in the area enclosed by f_1 and f_2 , these sources to have a rate of emission equal to $H(t)$ per unit area per unit time per side of the surface. Thus the potential at the point $Q(x, y, z)$ is given by

$$\Phi = 2 \iint_{\Delta} \Phi_S[(x - \xi), (y - \eta), z, t] d\xi d\eta, \quad (4.1)$$

taken over the area Δ of the aerofoil. A source can contribute to this integral only if its zone of influence (figure 1) contains the field point Q . This means that the integrand is non-zero only in a certain region on the aerofoil forward of the point Q , as shown in figure 8 for $z = 0$, and is infinite along the forward Mach cone.

The first stage of this integration is implicit in the forms given for a needle source, and it is fairly clear that equivalent to (4.1) is the integral

$$\Phi = 2 \int_{y_1}^{y_2} \Phi_N[\{x - f_1(\eta)\}, \{y - \eta\}, z, t] d\eta - 2 \int_{y_1}^{y_2} \Phi_N[\{x - f_2(\eta)\}, \{y - \eta\}, z, t] d\eta, \quad (4.2)$$

where y_1 and y_2 are the ordinates of the wing tips.

Since disturbances cannot be propagated forward of the Mach cones and $|df_2/dy| < \cot \mu$, the effect of the second integral on the potential over the aerofoil is necessarily zero, and it will henceforth be disregarded. Similarly, the pressure change in the region forward of the envelope of Mach cones from the trailing edge is

$$p = 2 \int_{y_1}^{y_2} p_N[\{x - f_1(\eta)\}, \{y - \eta\}, z, t] d\eta. \quad (4.3)$$

It would appear that the pressure might alternatively be written

$$p = 2 \iint_{\Delta} p_S[(x - \xi), (y - \eta), z, t] d\xi d\eta, \quad (4.4)$$

where p_S is the pressure corresponding to Φ_S , supposing it to exist. We shall use this form, and it will be found that some results can be derived very quickly in this way, though with questionable rigour. They can be checked in various ways—for example, from (4.3)—and are quite correct.

Equations (4.2) and (4.3) describe integrations ranging over elemental needle sources arranged parallel to the x -axis, with their vertices on the leading edge. At any given time t , the centres of the spheres B of the elemental sources are distributed along the locus

$$x = f_1(y) + Ut, z = 0,$$

and in general these spheres have a tube-shaped envelope. This envelope, which corresponds to one of the characteristics of the wave-equation, is a locus on which changes in the analytic form of the potential may occur. If the function has a 'kink', more loci of the same kind appear, springing from the 'kink' in the form of the point source itself. These remarks are illustrated by figure 2 which shows the section of these loci by the plane $z = 0$. The leading edges are marked A , and the other solid lines in the figure, e.g. B and C , are loci of the kind referred to. The chain-dotted lines are loci of branch-points of multi-valued functions occurring in the expressions for the potential, pressure, etc. No change in essential analytic form takes place across them,

but they are obtrusive in the work since superficially different results are obtained in regions separated by them. The significance of these remarks will be clearer after reading §7. For points between the lines A and B , the integrals for the potential and pressure are independent of the time. This means that the final steady flow is established in this region, and spreads downstream in the wake of the envelope of spheres. Ultimately, of course, it envelops the whole of any finite wing, and then the lift has attained its final value. If the trailing edge is straight and at right angles

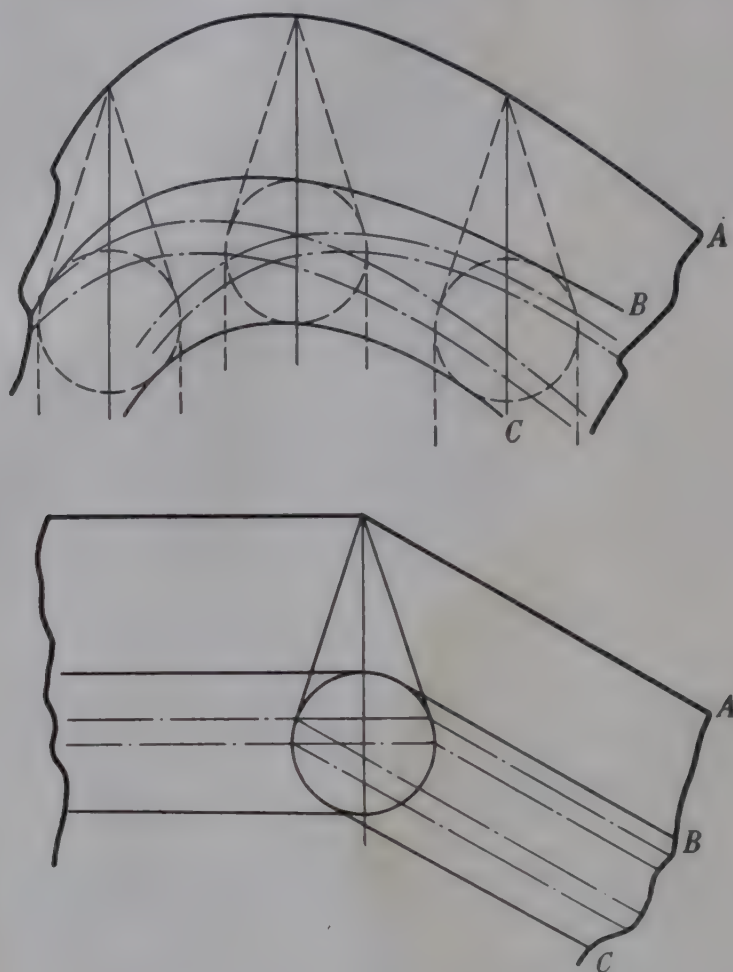


FIGURE 2. Sudden change of incidence.

to the direction of flight, this occurs when the greatest chord of the aerofoil has travelled $(1 - \sin \mu)^{-1}$ times its own length after the disturbance. It may also be observed that points behind the line C are not affected by the boundary A . All such points on the aerofoil will therefore have the same pressure and potential, which is independent of the plan-form. At the start of the motion, i.e. for $t \rightarrow 0$, the lines B and C coincide with A , and so the *initial pressure coefficient of a supersonic wing is constant and independent of the plan-form*. It is $\cos \mu$ times the Ackeret pressure for a two-dimensional wing in the same circumstances. The perturbation flow here is expansion normal to the wing, as if the source sheet were infinite in extent.

The plan-form shown in figure 3, which is the sector of the plane $z = 0$ bounded by the lines $y = 0$ and $y = x \tan \gamma > x \tan \mu$, is not a purely supersonic plan-form. However, a sheet of sources having such a shape is a valid subject for consideration, and *any* purely supersonic plan-form having a contour consisting of a finite number

of straight lines (for example, trapezoidal and delta wings) can be obtained from it by superposition. A plan-form like figure 3 will be called a 'triangular half-wing', and it is the most convenient unit for the preparation of charts. It is to be noted that for this particular shape, the form of the diagram does not change with time, but merely increases in size. We shall find that the pressures at corresponding points at different times are the same. Detailed consideration of this example follows later.

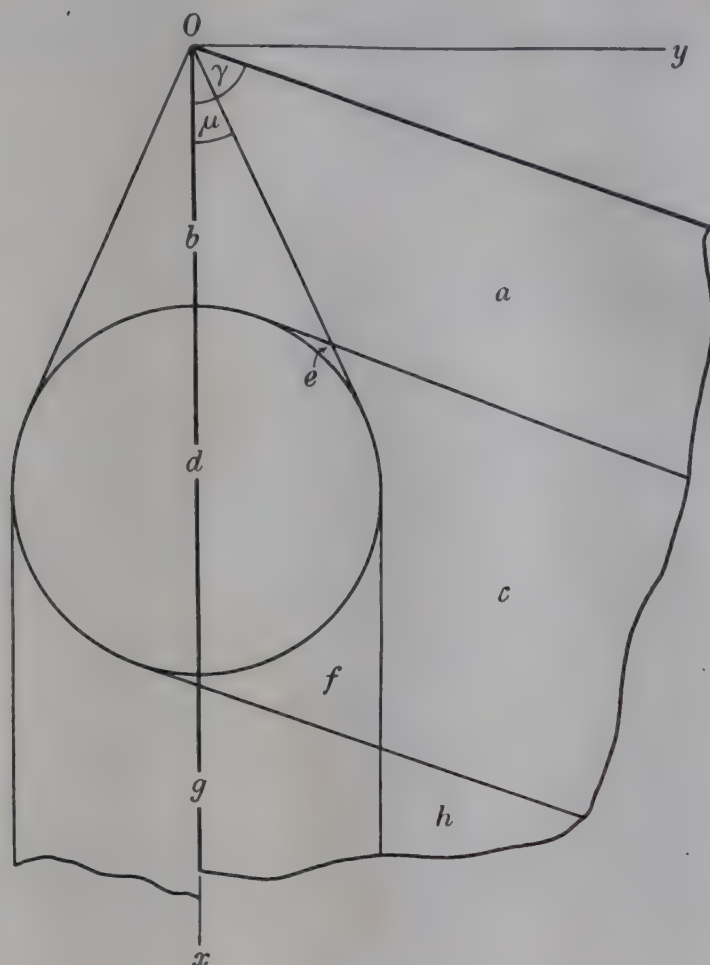


FIGURE 3. Triangular half-wing sudden change of incidence.

5. SHARP-EDGED GUST

As in § 4, the potential

$$\Phi = 2 \int_{y_1}^{y_2} \Phi_G[\{x - f_1(\eta)\}, \{y - \eta\}, z, \{t - f_1(\eta)/U\}] d\eta - \int_{y_1}^{y_2} \Phi_G[\{x - f_2(\eta)\}, \{y - \eta\}, z, \{t - f_2(\eta)/U\}] d\eta$$

represents an aerofoil entering a sharp-edged gust, and the second integral may be ignored. Similarly, the pressure change is

$$p = 2 \int_{y_1}^{y_2} p_G[\{x - f_1(\eta)\}, \{y - \eta\}, z, \{t - f_1(\eta)/U\}] d\eta. \quad (5.1)$$

At any time $t > 0$, the centres of the spheres B of the elemental gust sources are distributed along the line $x = Ut$, $z = 0$. Since their radius varies like $\{t - f_1(y)/U\}$,

their envelope is a tube which shrinks to a point at the edge of the aerofoil. Figure 4 shows the state of affairs on the plane $z = 0$. Between the leading edge A and the

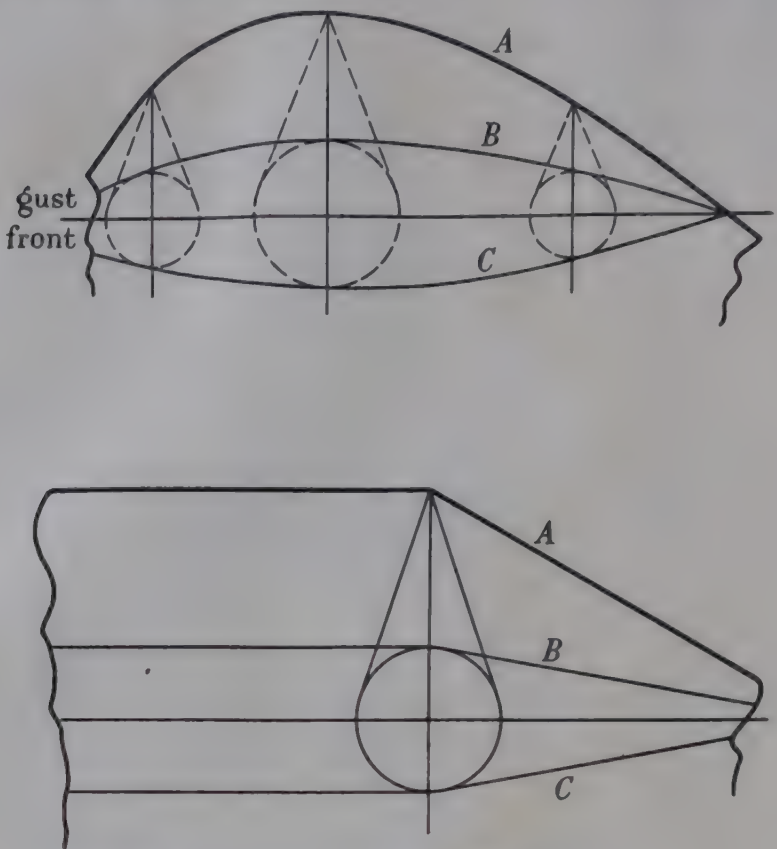


FIGURE 4. Entry into a sharp-edged gust.

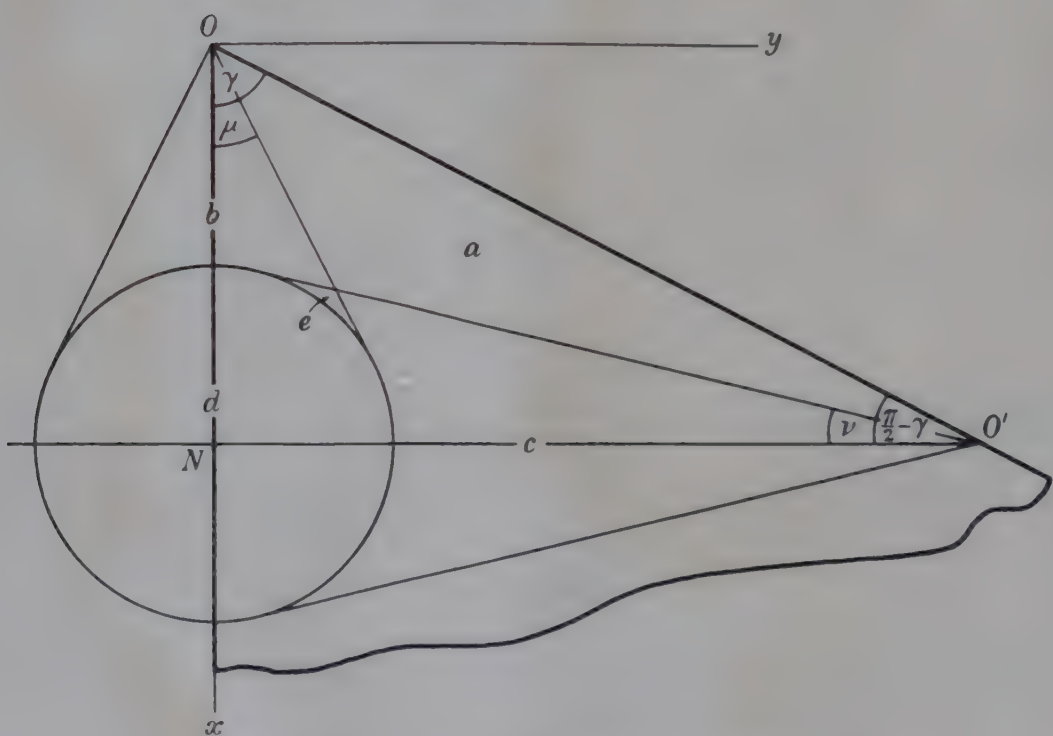


FIGURE 5. Triangular half-wing sharp-edged gust.

line B , the flow is the same as in the steady state, and behind the line C there is no disturbance at all. Figure 5 shows the important case of the triangular half-wing, which is treated in detail later.

6. THE CONICAL PRESSURE FUNCTION FOR A TRIANGULAR HALF-WING

For the rest of the paper, we shall take $z = 0$ and consider conditions over the aerofoil itself. In regions a, c, h of figure 3 the conditions are the same as for an infinite swept wing, because a point in one of these areas is not affected by the cut edge $y = z = 0$, and these again are the same as for an infinite unswept wing with velocity $U \sin \gamma$. The flow and pressure distributions in a, c and h are therefore known from the two-dimensional results, and in particular the pressure in a is

$$p_a = \rho U A \tan \mu,$$

where

$$A^2 = (1 - \tan^2 \mu / \tan^2 \gamma)^{-1}.$$

If we now define the pressure-growth function $\Psi_{1\alpha}$ as

$$\Psi_{1\alpha} = p / \rho U \tan \mu = \text{pressure/Ackeret pressure},$$

we have that

$$[\Psi_{1\alpha}]_a = A. \quad (6.1)$$

For regions a and b of figure 3 the potential ϕ_N in the integral of equation (4.2) is independent of the time, and it follows that in these regions the field of flow is the same as for a sheet of steady sources, which is fairly obvious on physical grounds. The pressure in b , obtained by evaluating (4.3), is

$$p_b = \frac{\rho U A \tan \mu}{\pi} \cos^{-1} \left\{ \frac{x \tan^2 \mu - y \tan \gamma}{(x \tan \gamma - y) \tan \mu} \right\},$$

so that

$$[\Psi_{1\alpha}]_b = \frac{A}{\pi} \cos^{-1} \left\{ \frac{x \tan^2 \mu - y \tan \gamma}{(x \tan \gamma - y) \tan \mu} \right\}.$$

This depends on the ratio y/x , which is to be expected since the flow is conical in the region under consideration. Let

$$\tan \phi = y/x,$$

and

$$P(\alpha, \beta) = \frac{1}{\pi} \cos^{-1} \left\{ \frac{\alpha - \beta}{1 - \alpha\beta} \right\},$$

so that

$$[\Psi_{1\alpha}]_b = AP \left\{ \frac{\tan \mu}{\tan \gamma}, \frac{\tan \phi}{\tan \mu} \right\}. \quad (6.2)$$

This function P , with different arguments α and β , has the property that all the pressures over the triangular half-wing (and plan-forms derived from it), for both unit cases considered, can be expressed in terms of it. Because the function occurs so often in computations, a chart of it is given as figure 6.

Relation to Puckett's formula

A complete delta wing is obtained by adding two half-wings. In the central Mach cone, the pressure is

$$p = \rho U A \tan \mu \cdot \{P(\alpha, \beta) + P(\alpha, -\beta)\},$$

where

$$\alpha = \tan \mu / \tan \gamma \quad \text{and} \quad \beta = \tan \phi / \tan \mu,$$

i.e.
$$p = \frac{\rho U A \tan \mu}{\pi} \cos^{-1} \left\{ \frac{2\alpha^2 - \alpha^2 \beta^2 - 1}{1 - \alpha^2 \beta^2} \right\},$$

and this is identical with Puckett's result (Puckett 1946) when the necessary changes in notation are made.

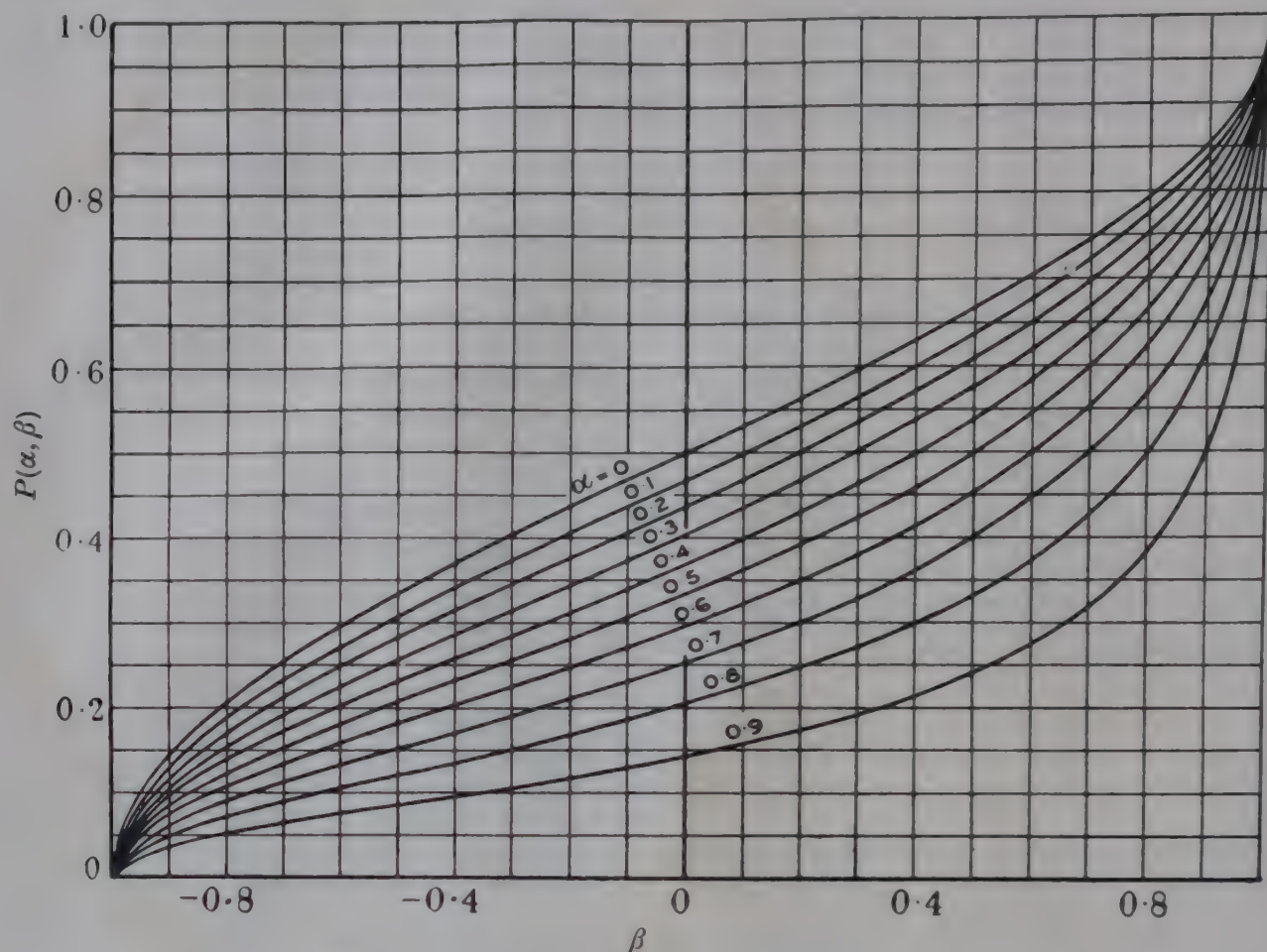


FIGURE 6. Conical pressure function $P(\alpha, \beta)$.

The edge of a sheet of sources

A sheet of sources of transient type covering the half-plane $y > 0$, $z = 0$ was considered in an earlier paper (Strang 1948), and the pressure change in the edge region was found to be

$$p = \frac{\rho V a}{4\pi} \{ \pi + 2 \sin^{-1} (y/at) \}.$$

This can be treated as a degenerate case of the triangular half-wing, in which $\mu \rightarrow 0$ so that $\alpha = 0$, $\beta = y/at$ and

$$p = \frac{1}{2} \rho V a_1 \cdot P(0, y/at),$$

which is equivalent to the former expression. As an application, the pressure-growth function in region g of figure 3 is

$$[\Psi_{1\alpha}]_g = \cos \mu \cdot P(0, y/at). \quad (6.3)$$

Pressure-growth functions for two-dimensional aerofoils

The pressure-growth functions Ψ_{1G} and $\Psi_{1\alpha}$ for two-dimensional aerofoils (Strang 1948) will be utilized later in this paper, and it is rather remarkable that they too

can be expressed compactly in terms of the function P . The expressions, which may be verified by direct substitution, are

$$\left. \begin{aligned} \Psi_{1G} &= 1 & \text{if } x < Ut(1 - \sin \mu), \\ \Psi_{1G} &= P\left(\sin \mu, \frac{Ut - x}{at}\right) & \text{if } Ut(1 - \sin \mu) < x < Ut(1 + \sin \mu), \\ \Psi_{1G} &= 0 & \text{if } x > Ut(1 + \sin \mu), \end{aligned} \right\} \quad (6.4)$$

and

$$\left. \begin{aligned} \Psi_{1\alpha} &= 1 & \text{if } x < Ut(1 - \sin \mu), \\ \Psi_{1\alpha} &= P\left(\sin \mu, \frac{Ut - x}{at}\right) + \cos \mu \cdot P\left(0, \frac{x - Ut}{at}\right) & \text{if } Ut(1 - \sin \mu) < x < Ut(1 + \sin \mu), \\ \Psi_{1\alpha} &= \cos \mu & \text{if } x > Ut(1 + \sin \mu). \end{aligned} \right\} \quad (6.5)$$

7. INCIDENCE CHANGE SOLUTIONS FOR A TRIANGULAR HALF-WING

The pressure distribution

The velocity potential and pressure distribution can be obtained by substitution of the appropriate limits in equations (4.2) and (4.3), and this is how the results were first obtained. A more instructive but less convincing method has since been devised, and this will be described first. The discussion will deal chiefly with the pressure distribution, which has the greater practical interest, but similar treatment of the velocity potential is possible.

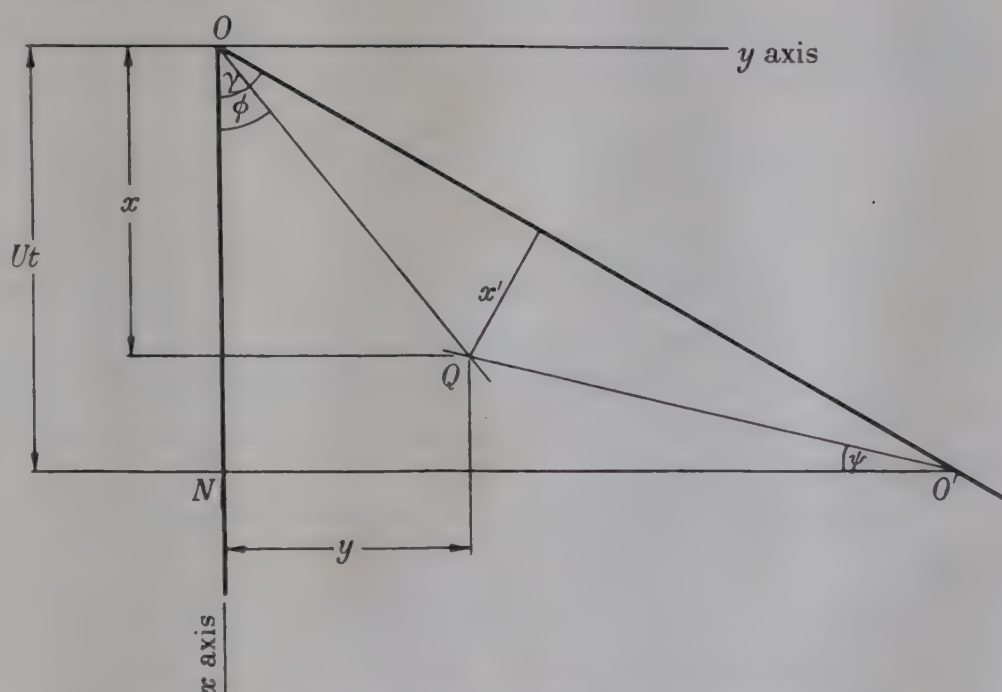


FIGURE 7. Scheme of co-ordinates for triangular half-wing.

We have found expressions for $[\Psi_{1\alpha}]_b$ and $[\Psi_{1\alpha}]_q$, and the observation that the flow in regions a , c and h behaves as if they were part of a two-dimensional unswept wing with velocity $U \sin \gamma$ enables us to write down at once the pressure in these regions. Let

$$x' = x \sin \gamma - y \cos \gamma,$$

so that x' is the distance of the point $Q(x, y)$ from the leading edge (figure 7). Making the appropriate changes in equation (6.5), we have that

$$\begin{aligned} [\psi_{1\alpha}]_a &= A, \\ [\psi_{1\alpha}]_c &= AP(\alpha', \beta') + \cos \mu P(0, -\beta'), \end{aligned} \quad (7.1)$$

$$[\psi_{1\alpha}]_h = \cos \mu, \quad (7.2)$$

where

$$\alpha' = \sin \mu / \sin \gamma,$$

$$\beta' = \frac{\sin \gamma}{\sin \mu} - \frac{x'}{at}.$$

Now consider the integral of equation (4.1) when $z = 0$. The integrand is essentially zero everywhere except in a certain region σ upstream from the field point $Q(x, y)$, this region having the shape of the field of influence of a supersonic point source. The effective field of integration is the area common to σ and Δ , shown shaded in figure 8. The form of the integral is such that

(i) *its value depends only on the shape of the field of integration and not on its position, and*

(ii) *for geometrically similar shapes, the integral is proportional to the scale.*

These results can be proved very easily by dimensional analysis. As an application of (i) we can deduce that lines parallel to the leading edge are equipotentials in a and c , and as an application of (ii) the potential in b is proportional to OQ along a line through O (figure 8).

The supersonic point source can scarcely be said to possess a pressure distribution, so curious is the potential. Let us, however, accept the existence of a discontinuous function expressing its pressure so that the integral

$$p = 2 \iint_{\Delta} p_S[(x - \xi), (y - \eta), 0, t] d\xi d\eta,$$

analogous to (4.1), is the pressure at the point $(x, y, 0)$. As for the potential, the effective field of integration is the area common to σ and Δ . Dimensional analysis shows that the value of the integral *depends only on the shape of the field of integration* and is independent of the size of the field, and its position. Let us apply this to the regions a to h of figure 3 to see what can be learned. Figure 8 should be studied in this connexion.

Region a. The field of integration has the same triangular shape for all points Q in a , so that the pressure is a constant.

Region b. The field of integration is a quadrilateral whose shape depends only on the angle ϕ , so that the pressure is constant on radial lines through O , i.e. a conical field.

Region c. The pressure is constant on lines parallel to the leading edge.

Region e. The field of integration for e can be obtained by superposition as $[c] - \{[a] - [b]\}$ so that

$$[\Psi_{1\alpha}]_e = [\Psi_{1\alpha}]_b + [\Psi_{1\alpha}]_c - [\Psi_{1\alpha}]_a.$$

Region f. The field of integration can be obtained by superposition as $[c] - \{[h] - [g]\}$ so that

$$[\Psi_{1\alpha}]_f = [\Psi_{1\alpha}]_c + [\Psi_{1\alpha}]_g - [\Psi_{1\alpha}]_h.$$

Region g. The pressure is constant on lines parallel to the x -axis.

Region h. The pressure is constant.

The pressures in a , b , c , g and h have been determined already (equations (6.1), (6.2), (7.1), (6.3) and (7.2)). This line of argument has enabled us to find the pressure distribution everywhere except in d , a defect the author has been unable to remedy.

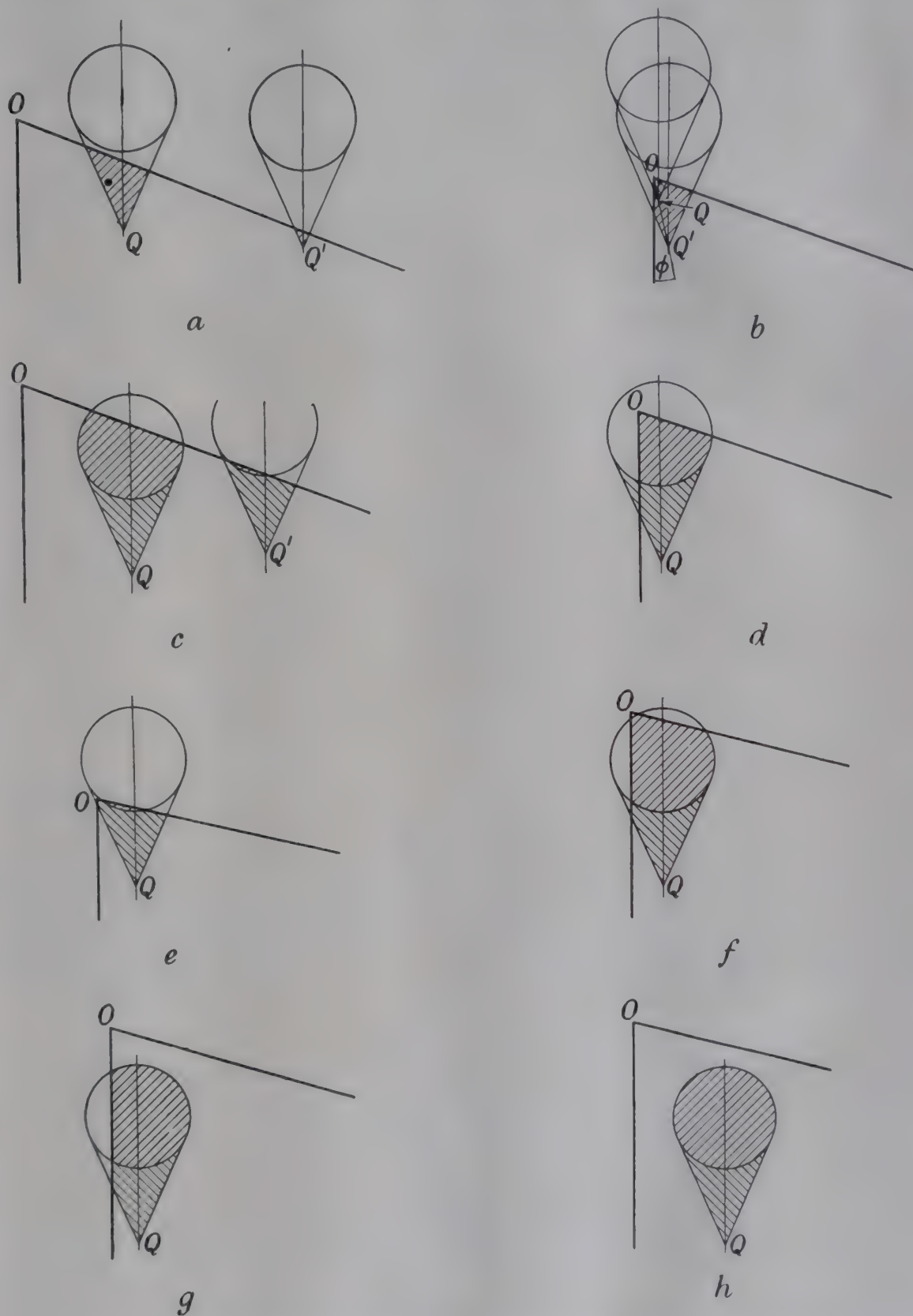


FIGURE 8. Fields of integration for regions a to h —change of incidence.

Some lack of confidence in the method may be felt, owing to the use of the pressure corresponding to a discontinuous potential.

Use of the formula (4.3) for the pressure is easier to justify, and was, in fact, the way the pressure distribution was first found. The integrations, and subsequent reduction to simple forms, present a rather formidable task, and I am indebted to Mr P. Wynter, of the University of Melbourne, for the performance of most of this work. The calculation of $[\Psi_{1\alpha}]_d$, necessary for the completion of the results, will be described to illustrate the procedure.

For the particular case of the triangular half-wing, the pressure integral becomes

$$p = 2 \int_0^\infty p_N[(x - \eta \cot \gamma), (y - \eta), 0, t] d\eta.$$

The integrand is conveniently broken into ranges with the end points y_1, y_2, y_3 , being the values of η for which the field point $Q(x, y)$ lies on:

- (i) the left edge of the Mach cone from the point $(y_1 \cot \gamma, y_1)$,
- (ii) the left edge of the circle of radius at with centre at $(Ut + y_2 \cot \gamma, y_2)$,
- (iii) the left edge of the cylindrical portion of the needle source.

As a matter of geometry, it may be verified that

$$\begin{aligned} y_1 &= \frac{x + y \cot \mu}{\cot \mu + \cot \gamma}, \\ y_2 &= y + (x' - Ut \sin \gamma) \cos \gamma + \sin \gamma \sqrt{a^2 t^2 - (x' - Ut \sin \gamma)^2}, \\ y_3 &= y + at. \end{aligned}$$

The integral for the pressure at Q may take any of the forms

$$p = \int_0^{y_1} p_B d\eta + \int_{y_1}^{y_2} p_A d\eta,$$

or

$$p = \int_0^{y_1} p_B d\eta,$$

or

$$p = \int_0^{y_1} p_B d\eta + \int_{y_2}^{y_3} p_C d\eta,$$

according to the position of Q within the circle, p_A, p_B and p_C being, of course, the pressures referred to in (3.3). These three forms give results which differ superficially, but can be reconciled by choosing the correct branches of the multi-valued inverse trigonometric functions which occur, and the following expression is valid throughout:

$$p = \frac{\rho UV}{4\pi} \{(\theta_1 + \theta_2) \sin \mu + A \tan \mu \cdot (\pi - \theta_3 + \theta_4)\},$$

where

$$\theta_1 = \sin^{-1} \left(\frac{y_2 - y}{at} \right) \quad (0 \leq \theta_1 \leq \pi \text{ increasing with } x'),$$

$$\theta_2 = \sin^{-1} (y/at) \quad (-\frac{1}{2}\pi \leq \theta_2 \leq \frac{1}{2}\pi),$$

$$\theta_3 = \sin^{-1} \left[\frac{(y_2 - y) \sin \gamma + A^2 x' \cot \gamma \tan^2 \mu}{A^2 x' \tan \mu} \right] \quad (0 \leq \theta_3 \leq \pi \text{ increasing with } x'),$$

$$\theta_4 = \sin^{-1} \left[\frac{y \sin \gamma - x \cos \gamma \tan^2 \mu}{x' \tan \mu} \right] \quad (-\frac{1}{2}\pi \leq \theta_4 \leq \frac{1}{2}\pi).$$

Putting $V = 2$, our result is that

$$[\Psi_{1\alpha}]_b = \frac{1}{2\pi} \{(\theta_1 + \theta_2) \cos \mu + A(\pi - \theta_3 + \theta_4)\},$$

and there is no obvious connexion between this and the results obtained for the other regions. The reduction to a comparable form will now be briefly sketched. We can show that

$$\begin{aligned} \frac{\theta_1}{\pi} &= \frac{1}{\pi} \cos^{-1} \left\{ \frac{Ut \sin \gamma - x'}{at} \right\} - \left(\frac{1}{2} - \frac{\gamma}{\pi} \right) \\ &= P(0, -\beta') - \left(\frac{1}{2} - \frac{\gamma}{\pi} \right), \end{aligned}$$

and
$$\begin{aligned} \frac{\theta_2}{\pi} &= \frac{1}{\pi} \cos^{-1} \left(-\frac{y}{at} \right) - \frac{1}{2} \\ &= P(0, y/at) - \frac{1}{2}, \end{aligned}$$

and
$$\begin{aligned} \frac{\theta_3}{\pi} &= \frac{1}{\pi} \cos^{-1} \left\{ \frac{Ut \sin^2 \mu - Ut \sin^2 \gamma + x' \sin \gamma}{x' \sin \mu} \right\} + \frac{1}{\pi} \cos^{-1} \left\{ \frac{\cos \gamma}{\cos \mu} \right\} - \frac{1}{2} \\ &= \frac{1}{2} - P(\alpha', \beta') + \frac{1}{\pi} \cos^{-1} (\cos \gamma / \cos \mu), \end{aligned}$$

and
$$\begin{aligned} \frac{\theta_4}{\pi} &= \frac{1}{2} - \frac{1}{\pi} \cos^{-1} \left\{ \frac{y \cot^2 \mu - x \cot \gamma}{\cot \mu (x - y \cot \gamma)} \right\} \\ &= \frac{1}{\pi} \cos^{-1} \left\{ \frac{x \cot \gamma - y \cot^2 \mu}{\cot \mu (x - y \cot \gamma)} \right\} - \frac{1}{2} \\ &= P \left\{ \frac{\tan \mu}{\tan \gamma}, \frac{\tan \phi}{\tan \mu} \right\} - \frac{1}{2}. \end{aligned}$$

Substituting these expressions,

$$\begin{aligned} 2[\Psi_{1\alpha}]_d &= AP(\alpha', \beta') + \cos \mu \cdot P(0, -\beta') + AP \left(\frac{\tan \mu}{\tan \gamma}, \frac{\tan \phi}{\tan \mu} \right) \\ &\quad + \cos \mu \cdot P(0, y/at) - \cos \mu + \frac{\gamma}{\pi} \cos \mu - \frac{A}{\pi} \cos^{-1} \left(\frac{\cos \gamma}{\cos \mu} \right) \\ &= [\Psi_{1\alpha}]_b + [\Psi_{1\alpha}]_c + [\Psi_{1\alpha}]_g - [\Psi_{1\alpha}]_h + \frac{\gamma}{\pi} \cos \mu + \frac{A}{\pi} \cos^{-1} \left(\frac{\cos \gamma}{\cos \mu} \right). \end{aligned}$$

The analysis for the other regions is similar, and the results agree with those stated earlier. It is convenient to introduce the non-dimensional variables

$$\begin{aligned} s &= Ut / (\text{chord of aerofoil}), \\ X &= x / (\text{chord of aerofoil}), \\ Y &= y / (\text{chord of aerofoil}), \\ X' &= X \sin \gamma - Y \cos \gamma, \end{aligned}$$

and gather up the results (table 1). As one would expect, the pressure is continuous across all internal boundaries.

TABLE 1. PRESSURE DISTRIBUTION FOR TRIANGULAR HALF-WING.
SUDDEN CHANGE OF INCIDENCE

region	$\Psi_{1\alpha} = \frac{\text{pressure}}{\text{Ackeret pressure}}$
<i>a</i>	<i>A</i>
<i>b</i>	<i>AP</i> (α, β)
<i>c</i>	<i>AP</i> (α', β') + $\cos \mu P(0, -\beta')$
<i>d</i>	$\frac{1}{2}([\Psi_{1\alpha}]_b + [\Psi_{1\alpha}]_c + [\Psi_{1\alpha}]_e - [\Psi_{1\alpha}]_h) + \frac{\gamma}{2\pi} \cos \mu - \frac{A}{2\pi} \cos^{-1}(\cos \gamma / \cos \mu)$
<i>e</i>	$[\Psi_{1\alpha}]_b + [\Psi_{1\alpha}]_c - [\Psi_{1\alpha}]_a$
<i>f</i>	$[\Psi_{1\alpha}]_c + [\Psi_{1\alpha}]_e - [\Psi_{1\alpha}]_h$
<i>g</i>	$\cos \mu P\{0, (Y/s) \operatorname{cosec} \mu\}$
<i>h</i>	$\cos \mu$

$X = x/Ut, \quad Y = y/Ut, \quad X' = X \sin \gamma - Y \cos \gamma, \quad s = Ut/c.$
 $\alpha = \tan \mu / \tan \gamma, \quad \beta = \tan \phi / \tan \mu,$
 $\alpha' = \sin \mu / \sin \gamma, \quad \beta' = \sin \gamma / \sin \mu - (X'/s) \sin \mu.$
 $P(x, y) = \frac{1}{\pi} \cos^{-1} \left\{ \frac{x-y}{1-xy} \right\} \quad (0 \leq P \leq 1).$
 $A^2 = (1 - \alpha^2)^{-1}.$

The lift- and moment-growth functions

In applications, combinations of triangular half-wings must be used such that the disturbed region to the left of the *x*-axis falls on the wing. We are, therefore, interested in the lift

$$2 \int_0^c dx \int p dy,$$

where the integration with respect to *y* covers the whole disturbance. This is, for example, half the lift of a complete delta wing. It is convenient to define a lift-growth function,

$$\Psi_{2\alpha} = \frac{\text{lift}}{\text{final lift}} = \frac{2}{\tan \gamma} \int_0^1 dX \int \psi_{1\alpha} dY,$$

which is directly applicable to the half-wing, in the sense explained, or to a complete symmetrical delta. It is a function of γ, μ and *s*. In the case of more general plan-forms built up from triangular half-wing elements, superposition may be used. Consider the plan-form shown in figure 9 and let

$$S_1 = \frac{1}{2} c_1^2 \tan \gamma_1, \quad S_2 = \frac{1}{2} c_2^2 \tan \gamma_2, \quad S_3 = \frac{1}{2} c_2^2 \tan \gamma_1, \quad S = S_1 + S_2 - S_3,$$
$$s_1 = Ut/c_1, \quad s_2 = Ut/c_2.$$

Then the lift-growth function for the whole wing is

$$\Psi_{2\alpha} = \frac{S_1}{S} \Psi_{2\alpha}(\gamma_1, s_1) + \frac{S_2}{S} \Psi_{2\alpha}(\gamma_2, s_2) - \frac{S_3}{S} \Psi_{2\alpha}(\gamma_1, s_2),$$

obtained by weighting the functions for individual pieces in the ratios of their areas. The Mach angle μ is the same for all, of course.

In a similar way, we define the moment-growth function for pitching moment about the leading edge as

$$\Psi_{3\alpha} = \frac{\text{leading edge pitching moment}}{\text{final pitching moment}} = \frac{3}{\tan \gamma} \int_2^1 X dX \int \Psi_{1\alpha} dY,$$

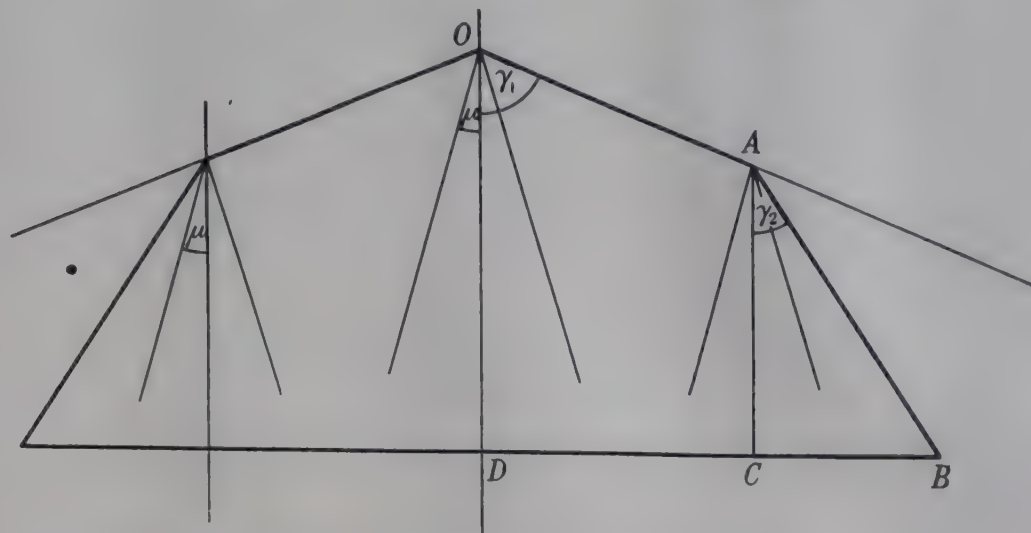


FIGURE 9. Illustrative supersonic wing plan-form $OD = c_1$, $AC = c_2$.

where, as before, the integration with respect to Y covers the whole surface. For the wing of figure 9

$$\Psi_{3\alpha} = \frac{2c_1 S_1}{M} \psi_{3\alpha}(\gamma_1, s_1) + \frac{1}{M} (3c_1 - c_2) \{S_2 \psi_{3\alpha}(\gamma_2, s_2) - S_3 \psi_{3\alpha}(\gamma_1, s_2)\},$$

where

$$M = 2cS_1 + (3c - c_2)(S_2 - S_3).$$

The functions $\Psi_{2\alpha}$ and $\Psi_{3\alpha}$ for the triangular half-wing have been calculated for three values of γ and for three values of the ratio $\tan \mu / \tan \gamma$, and the results are plotted in figures 10 to 15.

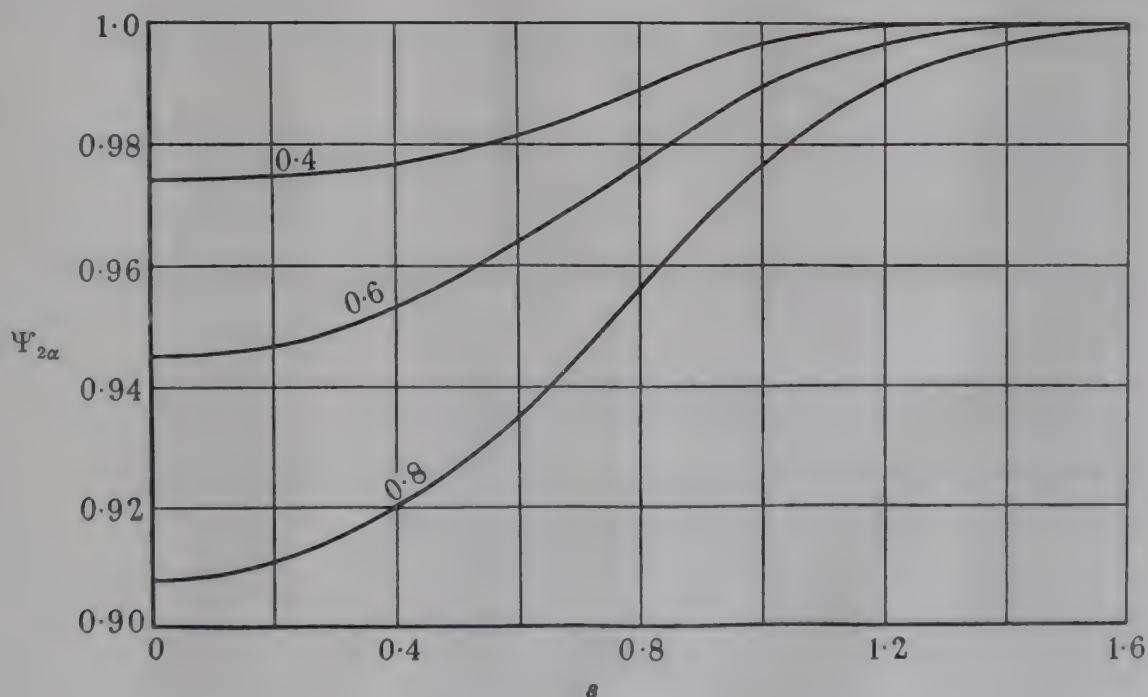


FIGURE 10. $\gamma = 30^\circ$.

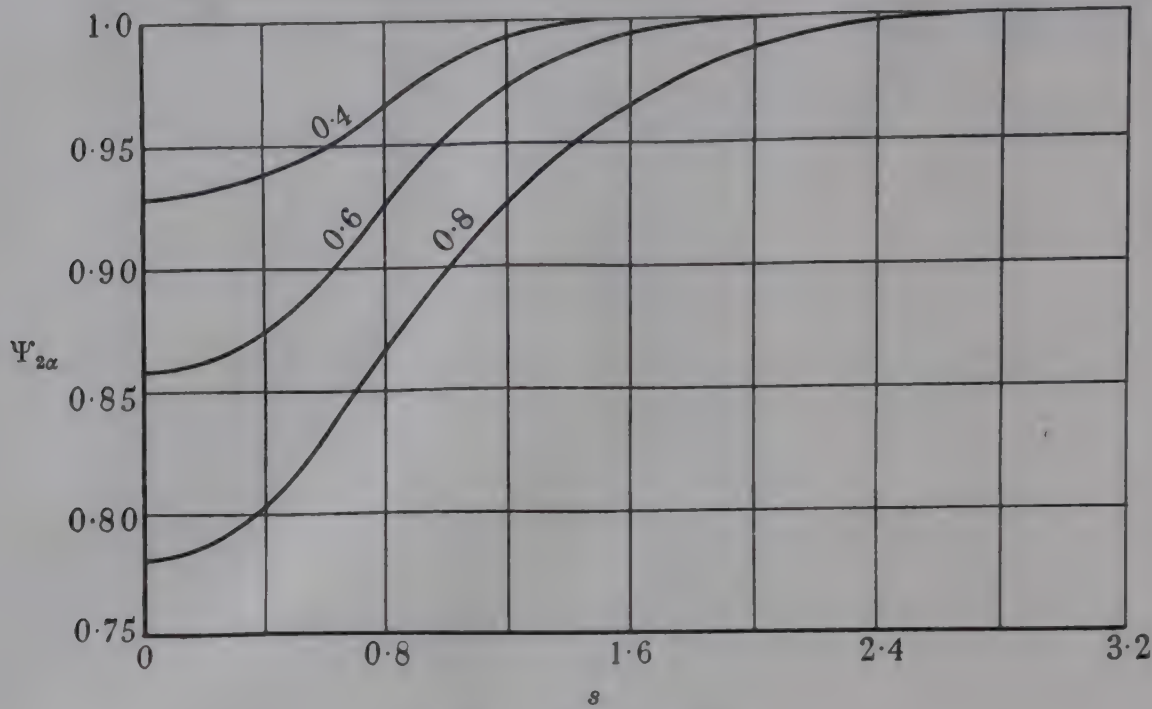


FIGURE 11. $\gamma = 45^\circ$.

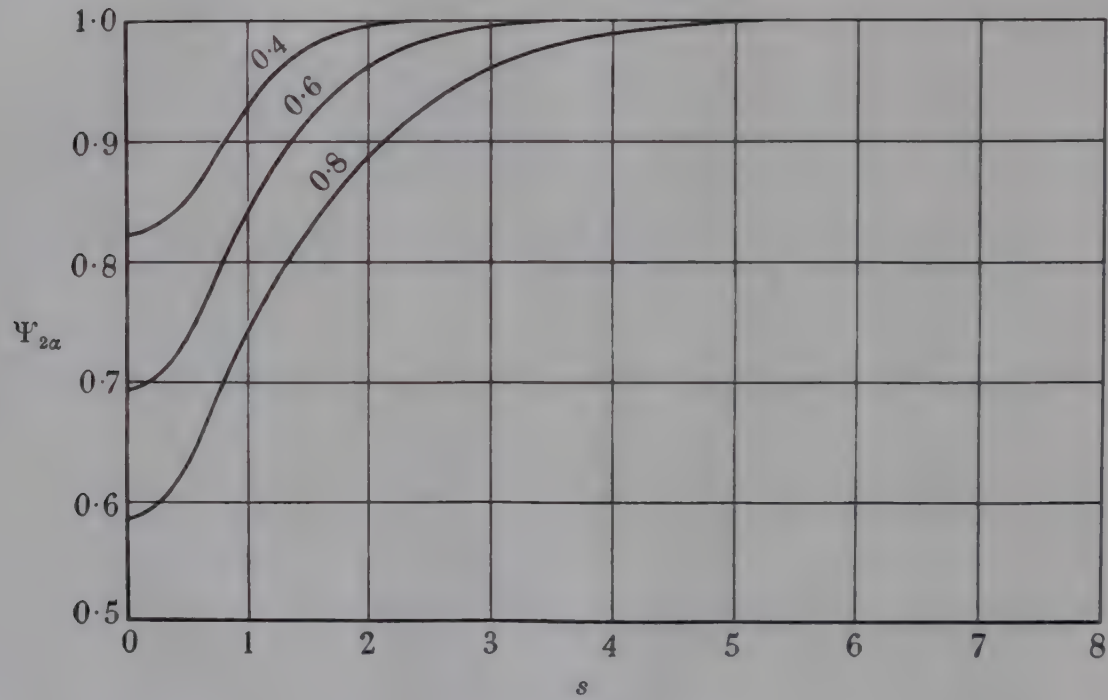


FIGURE 12. $\gamma = 60^\circ$.

FIGURES 10 TO 12. Lift growth function for change of incidence. s = chords travelled.

$$\Psi_{2\alpha} = \frac{\text{lift due to change of incidence}}{\text{Ackeret lift}}.$$

The numbers on the curves give the values of $\tan \mu / \tan \gamma$.

The full value of the lift is obtained when the wing has travelled $1/(1 - \sin \mu)$ chords, and this result is applicable to all purely supersonic wings that have straight trailing edges. The initial value of the lift is $\cos \mu$ times the Ackeret value, and this is true for *all* plan-forms, purely supersonic or otherwise. There is a curious consequence in the case of wings with very large sweep-back. The final lift of such wings may be less than the initial value, and the implications of this have not yet been worked out. It would seem that such plan-forms are extremely advantageous from the point of view of gust loading, which may well be a matter of importance. For purely supersonic delta wings, the dependence of $\Psi'_{2\alpha}$ on γ is not marked, the effect

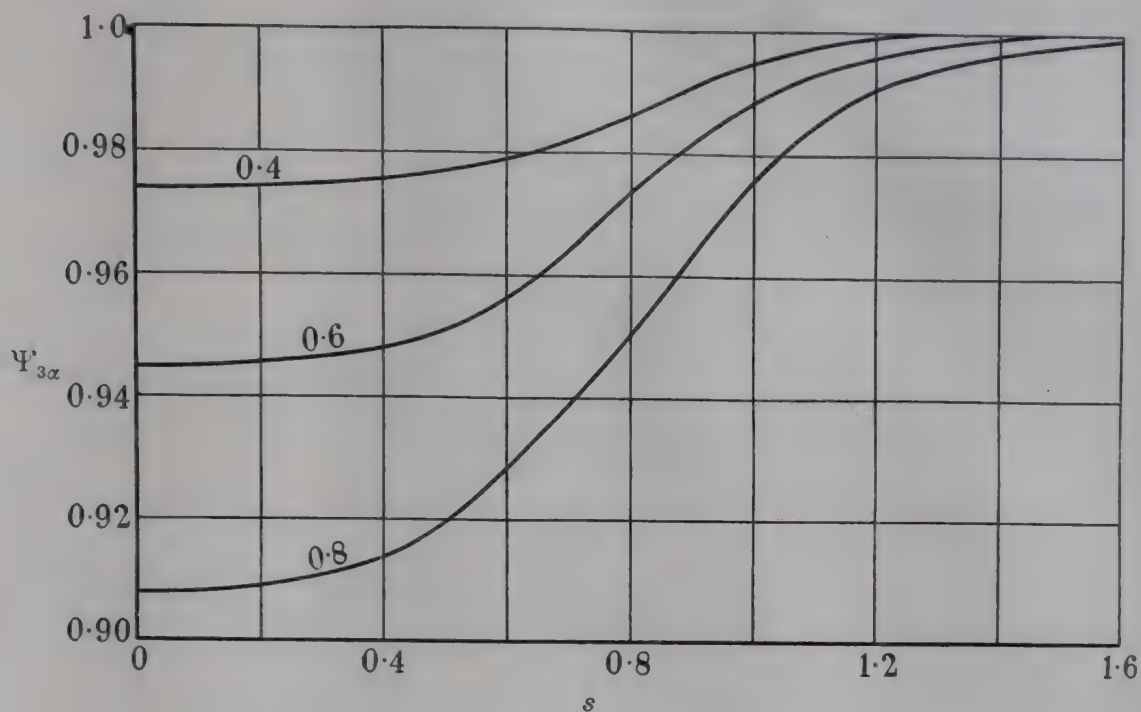


FIGURE 13. $\gamma = 30^\circ$.

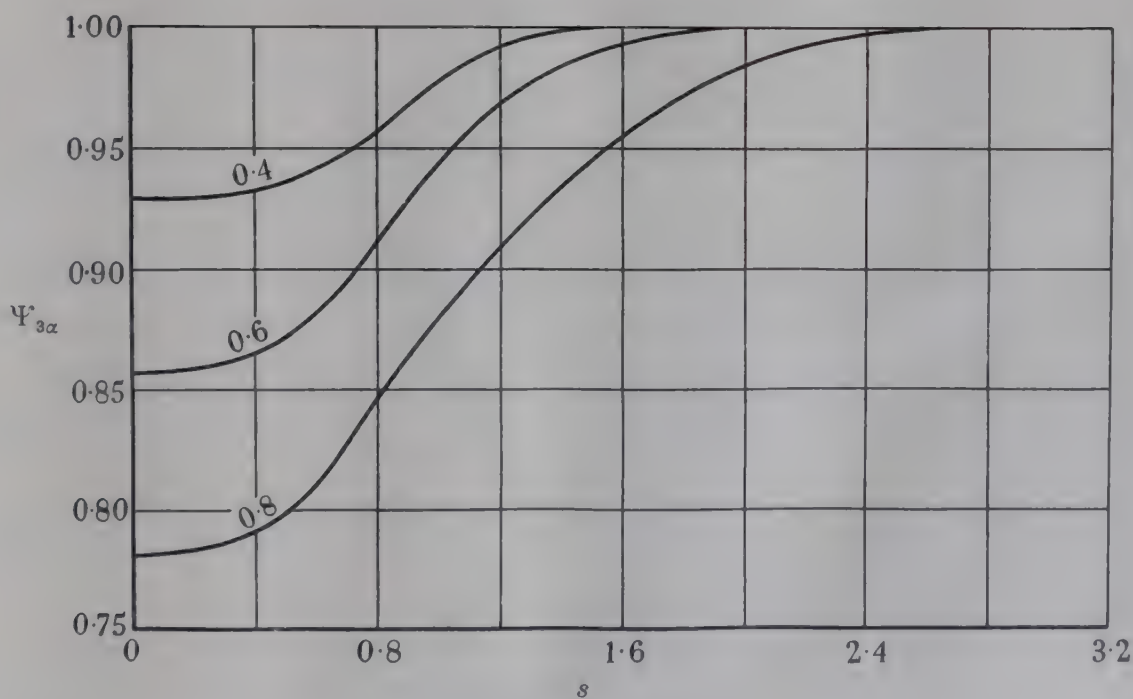
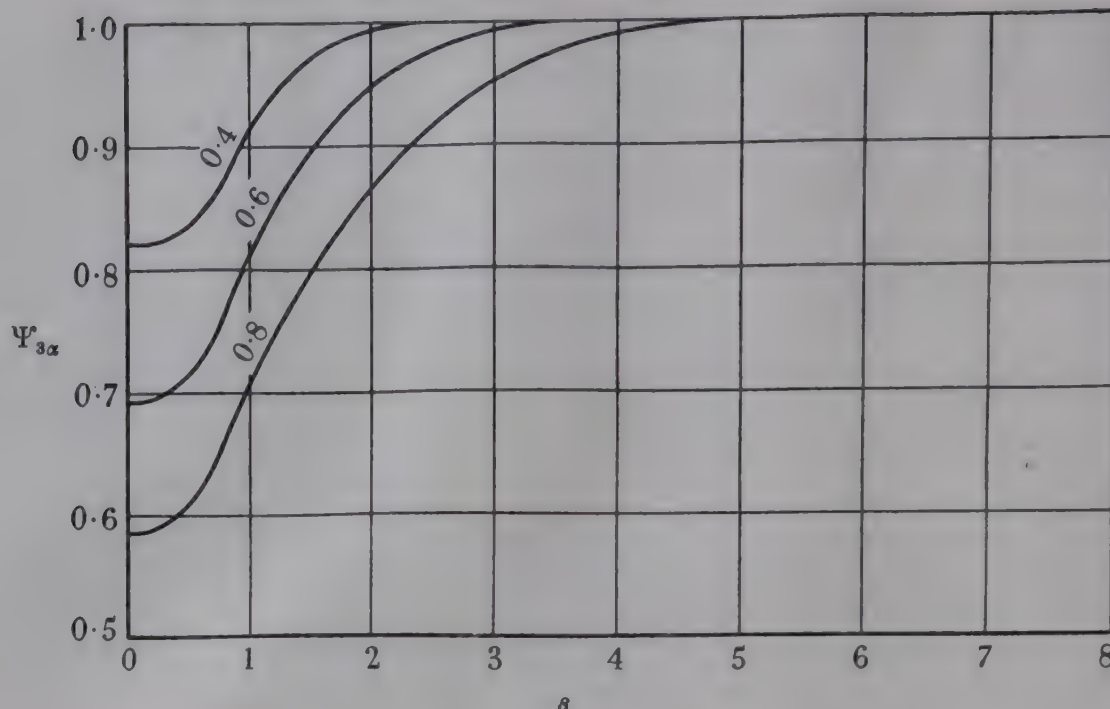


FIGURE 14. $\gamma = 45^\circ$.

FIGURE 15. $\gamma = 60^\circ$.

FIGURES 13 TO 15. Moment growth function for change of incidence. s = chords travelled.

$$\Psi_{3\alpha} = \frac{\text{leading edge moment due to change of incidence}}{\text{final leading edge moment}}.$$

The numbers on the curves give the values of $\tan \mu / \tan \gamma$.

of μ being much the more important. This implies that if two delta wings of equal area are made with different sweep-back, the variation in lift-growth is mainly due to the change in chord.

8. SHARP-EDGED GUST SOLUTION FOR TRIANGULAR HALF-WINGS

The pressure distribution

Viewed from a frame of reference fixed in the fluid, the problem is that of determining the pressure change produced by a triangular sheet of sources contained in the expanding triangle ONO' (figure 5), bounded by the fixed lines ON , $O'N$ and the moving line OO' . This may be regarded as a representation of a wing moving in the x -direction, with a gust front $O'N$, or equally as representing a wing moving in the y -direction with a gust front ON . We anticipate, therefore, a certain symmetry in the expressions for the potential, pressure, etc. To give expression to this, define a field point Q in the plane $z = 0$ by the bi-polar co-ordinates (ϕ, ψ) as in figure 5, and introduce the angle ν to correspond with μ . The parameter corresponding to $\tan \mu / \tan \gamma$ is $\tan \nu / \cot \gamma$, and figure 16, giving the relation between them, is useful for computing.

The pressure in regions a and b of figure 5 is, of course, the final steady pressure, so we can write down at once

$$\begin{aligned} [\Psi_{1G}]_a &= [\Psi_{1\alpha}]_a = A, \\ [\Psi_{1G}]_b &= [\Psi_{1\alpha}]_b = AP \left(\frac{\tan \mu}{\tan \gamma}, \frac{\tan \phi}{\tan \mu} \right). \end{aligned}$$

Using the symmetry, we can add the result

$$[\Psi_{1G}]_c = AP \left(\frac{\tan \nu}{\cot \gamma}, \frac{\tan \psi}{\tan \nu} \right).$$

Region e comes under the influence of two edges, and it is easily argued that

$$\begin{aligned} [\Psi_{1G}]_e &= [\Psi_{1G}]_a - \{[\Psi_{1G}]_a - [\Psi_{1G}]_b\} - \{[\Psi_{1G}]_a - [\Psi_{1G}]_c\} \\ &= [\Psi_{1G}]_b + [\Psi_{1G}]_c - [\Psi_{1G}]_a. \end{aligned}$$

This leaves region d unaccounted for, and once again this is more difficult. We revert to equation (5.1) which is applicable to all the wing, and in this case it becomes

$$p = 2 \int p_G \left[(x - \eta \cot \gamma), (y - \eta), 0, \left(t - \frac{\eta \cot \gamma}{U} \right) \right] d\eta.$$

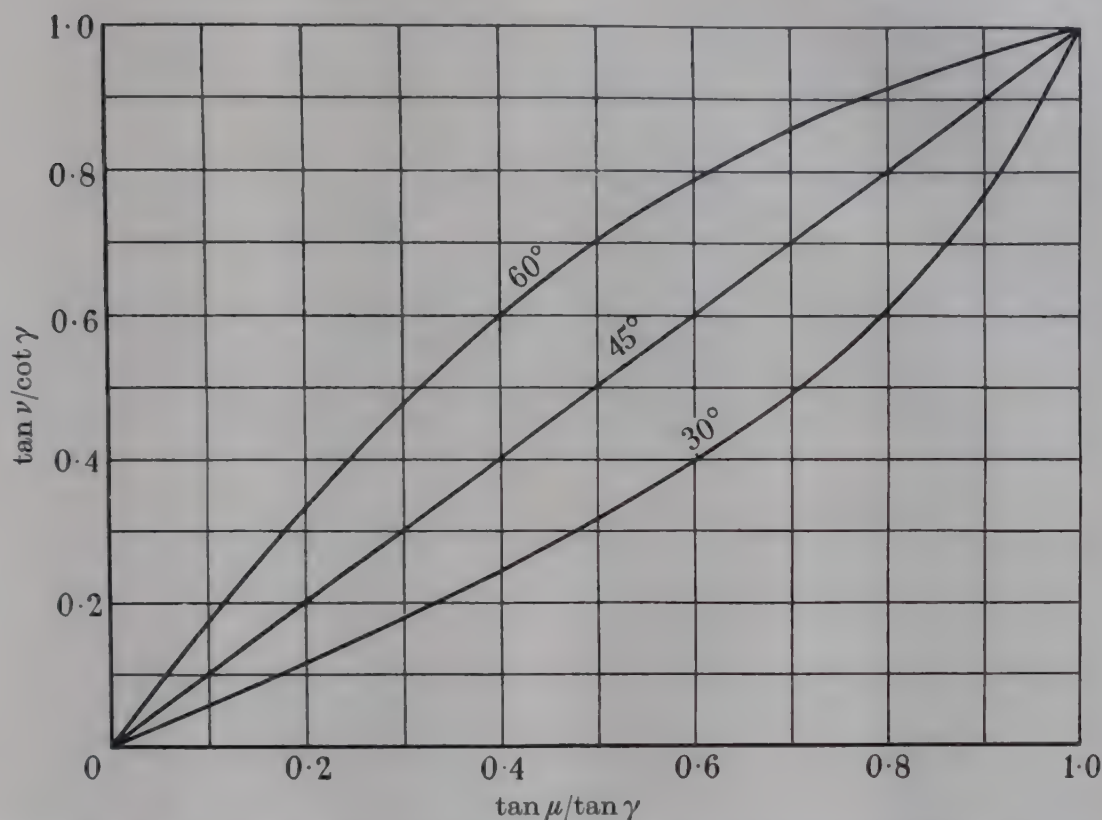


FIGURE 16. Curves for the determination of ν . The values of γ are indicated on the curves.

TABLE 2. PRESSURE DISTRIBUTION FOR TRIANGULAR HALF-WING.
ENTRY INTO A SHARP-EDGED GUST

region	$\Psi_{1G} = \frac{\text{pressure}}{\text{Ackeret pressure}}$
a	A
b	$AP(\alpha, \beta)$
c	$AP(\zeta, \eta)$
d	$\frac{A}{2} \left\{ [\Psi_{1G}]_b + [\Psi_{1G}]_c - \frac{1}{\pi} \cos^{-1}(\tan \mu \tan \nu) \right\}$
e	$[\Psi_{1G}]_b + [\Psi_{1G}]_c - [\Psi_{1G}]_a$
$\zeta = \tan \nu / \cot \gamma, \quad \eta = \tan \psi / \tan \nu.$	

The evaluation and reduction of this integral to simple form is very similar to that discussed in the last section, and will not be repeated. For a , b , c and e it yields the results already quoted, and for d it gives

$$[\Psi_{1G}]_d = \frac{A}{2} \left\{ P\left(\frac{\tan \mu}{\tan \gamma}, \frac{\tan \phi}{\tan \mu}\right) + P\left(\frac{\tan \nu}{\cot \gamma}, \frac{\tan \psi}{\tan \nu}\right) - \frac{1}{\pi} \cos^{-1}(\tan \mu \tan \nu) \right\},$$

or
$$2[\Psi_{1G}]_d = [\Psi_{1G}]_b + [\Psi_{1G}]_c - \frac{A}{\pi} \cos^{-1}(\tan \mu \tan \nu).$$

The lift- and moment-growth functions

These are defined as

$$\Psi_{2G} = \frac{\text{lift due to gust}}{\text{final lift}} = \frac{2}{\tan \gamma} \int_0^1 dX \int \Psi_{1G} dY$$

$$\Psi_{3G} = \frac{\text{leading edge moment due to gust}}{\text{final moment}} = \frac{3}{\tan \gamma} \int_0^1 X dX \int \Psi_{1G} dY,$$

and the calculated values for triangular half-wings are plotted in figures 17 to 22. The method of superposition for more general plan-forms, such as that shown in figure 9, differs slightly from the incidence-change case, because allowance must be made for the penetration into the gust. This allowance is made by re-defining

$$s_1 = Ut/c_1, \quad s_2 = \{Ut - (c_1 - c_2)\}/c_2 = 1 - \frac{c_1}{c_2} + \frac{Ut}{c_2},$$

and then the same formulae apply.

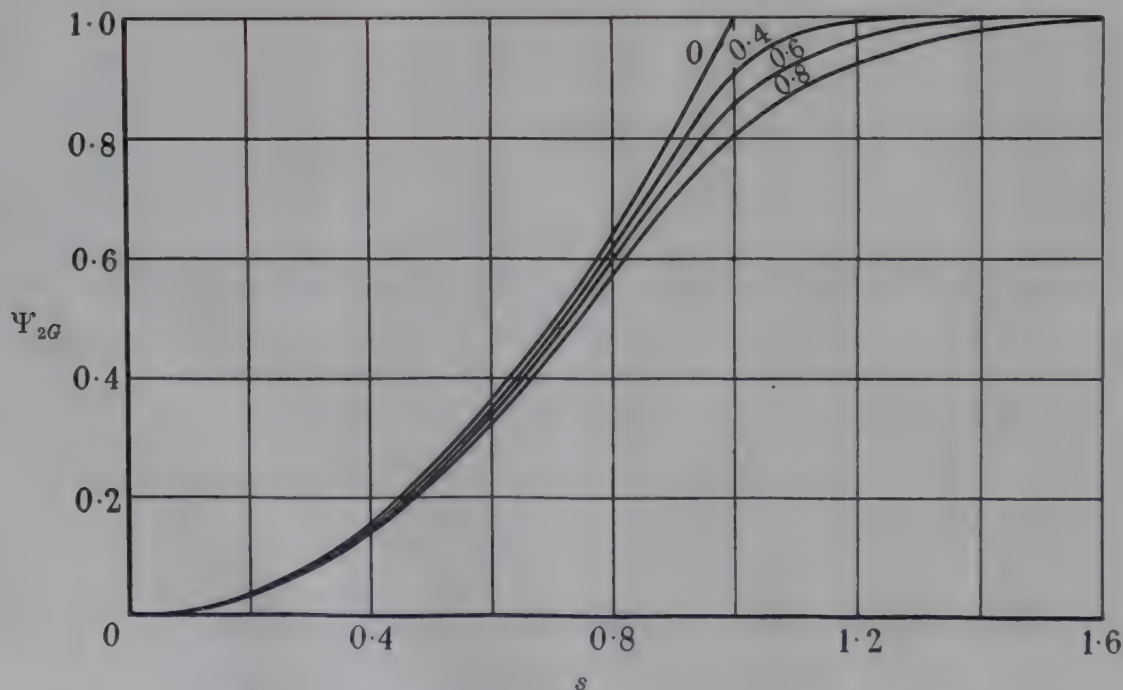


FIGURE 17. $\gamma = 30^\circ$.

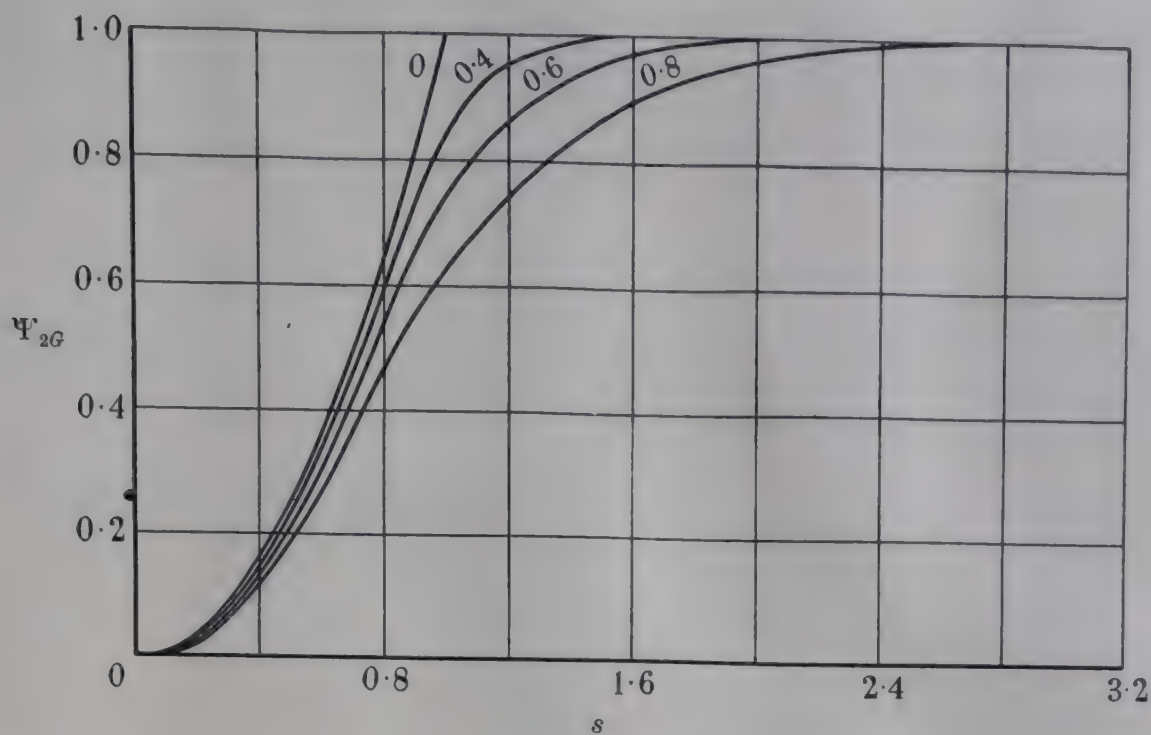


FIGURE 18. $\gamma = 45^\circ$.

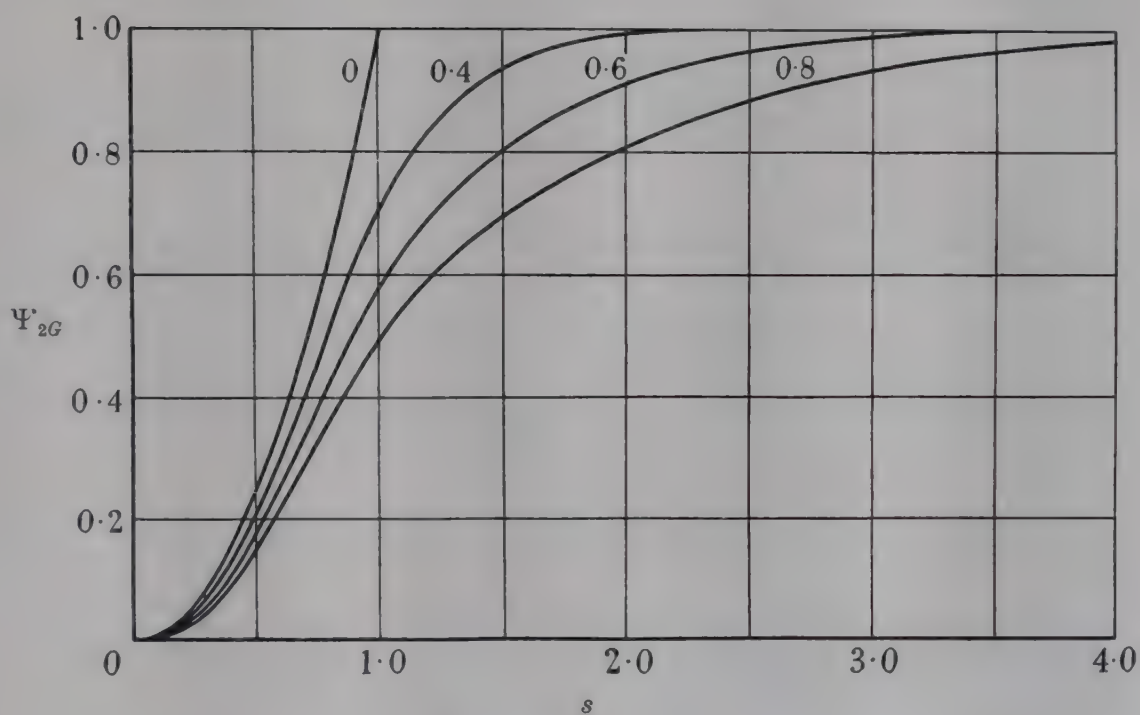


FIGURE 19. $\gamma = 60^\circ$.

FIGURES 17 TO 19. Lift-growth function for sharp-edged gust. s = chords travelled.

$$\Psi_{2g} = \frac{\text{lift due to gust}}{\text{steady lift}}.$$

The numbers on the curves give the values of $\tan \mu / \tan \gamma$.

The lift starts from zero at the instant of entry into the gust, and reaches its final value when a distance of $1/(1 - \sin \mu)$ chords has been traversed. This rule is applicable to any purely supersonic wing with a straight trailing edge. In the case of the delta wings, it is found that Ψ_{2G} is very little dependent on γ at Mach numbers greater than 2.

9. THE CASE $\gamma = 90^\circ$

The expressions for $\Psi_{1\alpha}$ and Ψ_{1G} retain their validity when $\gamma = 90^\circ$, and naturally simplify somewhat. But it is *not* true that the lift- and moment-growth functions tend to the two-dimensional ones as $\gamma \rightarrow 90^\circ$.

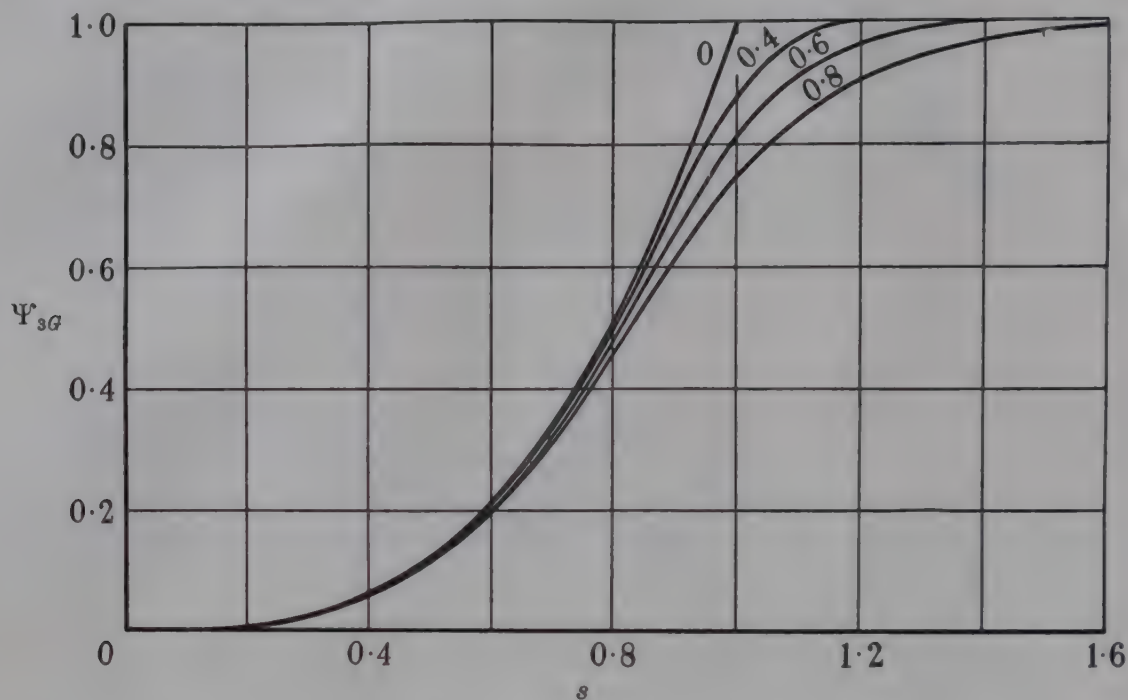


FIGURE 20. $\gamma = 30^\circ$.

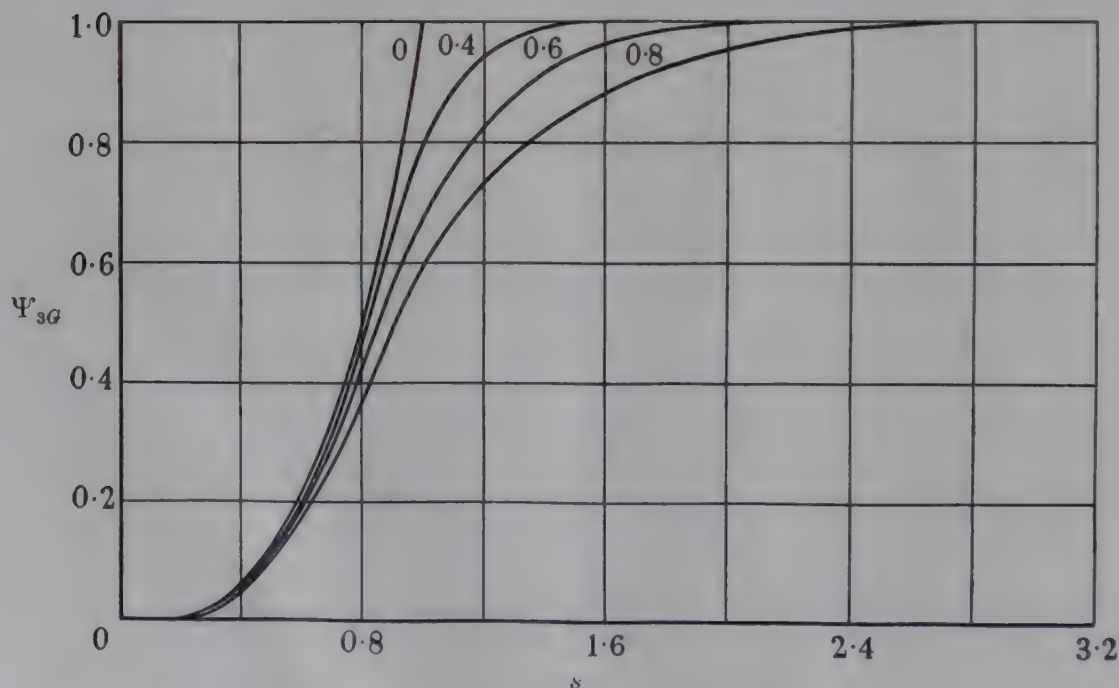


FIGURE 21. $\gamma = 45^\circ$.

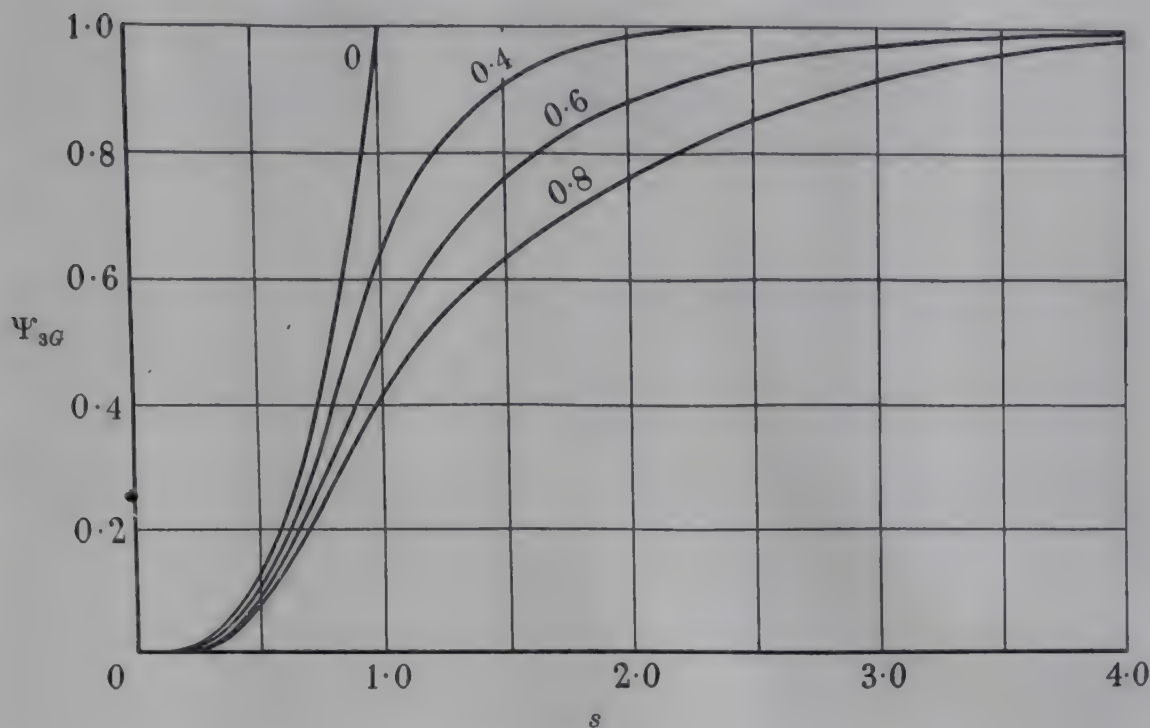


FIGURE 22. $\gamma = 30^\circ$.

FIGURES 20 TO 22. Moment growth function for sharp-edged gust. s = chords travelled.

$$\Psi_{3G} = \frac{\text{leading edge moment due to gust}}{\text{final leading edge moment}}.$$

The numbers on the curves give the values of $\tan \mu / \tan \gamma$.

This remark requires amplification. To illustrate the matter, consider a triangle bounded by the lines $y = 0$, $y = x \tan \gamma$, $x = c$, and suppose that the pressure is a function of the distance from the leading edge only, i.e. $p = f(x'/c)$. This is a simplification which is nearly true for large values of γ . Then it is easy to show that the mean pressure is

$$2 \int_0^1 f(\chi \sin \gamma) (1 - \chi) d\chi,$$

where $\chi = x'/c \sin \gamma$. As $\gamma \rightarrow 90^\circ$, this tends to

$$2 \int_0^1 f(\chi) (1 - \chi) d\chi.$$

For the two-dimensional strip $y > 0$, $0 < x < c$, on the other hand, the mean pressure is

$$\int_0^1 f(\chi) d\chi,$$

and these are only equal for the limited class of functions for which

$$\int_0^1 f(\chi) d\chi = 2 \int_0^1 f(\chi) \chi d\chi,$$

of which the most obvious example is $f(\chi) = \text{constant}$. Even in this case, the centres of pressure are quite different, being on the line $x = \frac{2}{3}c$ in the first case and $x = \frac{1}{2}c$ in the second.

In practice the difficulty is readily overcome. Parallel portions of wings always terminate finitely at both ends, and the lift- and moment-growth functions for such a piece of wing, *including the disturbed tip regions*, are the two-dimensional ones.

CONCLUSIONS

For any purely supersonic aerofoil experiencing a sudden change of incidence, or entering a sharp-edged gust, the pressure distribution at later times is calculable, and specific results are given in this paper for plan-forms made up of straight-line segments. Charts are provided for the determination of the lift and moment in such cases. Generalization to arbitrary motions (without pitching) and to arbitrary gusts is by means of the superposition integral.

The findings of special interest are

- (i) the properties of the function P , and
- (ii) the fact that the lift coefficient at any time depends mainly on the Mach number and the chord-distance travelled, and little on the sweep-back. This implies that plan-forms of small aspect-ratio are the most advantageous from the point of view of gust loading, on account of their greater chord-lengths for given area.

This work was carried out at the Aeronautical Research Laboratories, Fishermen's Bend, Melbourne, and is published by permission of the Commonwealth Department of Supply and Development, Australia.

REFERENCES

- Bratt, J. B. & Chinneck, A. 1947 *A.R.C. Rep.* no. 10,709.
 Evvard, J. C. 1948 *N.A.C.A. Tech. Note*, no. 1699.
 Garrick, I. E. & Rubinow, S. I. 1947 *N.A.C.A. Tech. Note*, no. 1383.
 Heaslet, N. A. & Lomax, H. 1947 *N.A.C.A. Tech. Note*, no. 1515.
 Jones, W. P. 1947 *A.R.C. Rep.* no. 10,871.
 Possio, C. 1937 *Acta Pontif., Acad. Sci.* 1. Translated as *A.R.C. Rep.* no. 7668.
 Puckett, A. E. 1946 *J. Aero. Sci.* 13, no. 9.
 Robinson, A. 1947 *Coll. Aeronaut. Cranfield, Rep.* no. 9.
 Robinson, A. 1948 *Coll. Aeronaut. Cranfield, Rep.* no. 16.
 Schwarz, L. 1943 *Jb. Lufo.* 1943 I.A. 010. Translated as *M.A.P.* -VG59-188T.
 Strang, W. J. 1948 *Proc. Roy. Soc. A*, 195, 245.
 Strang, W. J. 1950 *Proc. Roy. Soc. A*, 202, 40.

The instability of liquid surfaces when accelerated in a direction perpendicular to their planes. II

By D. J. LEWIS, *Engineering Laboratory, University of Cambridge*

(Communicated by Sir Geoffrey Taylor, F.R.S.—Received 3 December 1949)

[Plates 5 to 11]

An apparatus for accelerating small quantities of various liquids vertically downwards at accelerations of the order of $50g$ (g being 32.2 ft./sec.²) is described, and the behaviour of small wave-like corrugations initially imposed on the upper liquid surface has been observed by means of high-speed shadow photography. The instability observed under a wide variety of experimental conditions has been analyzed, and the initial phases have been found to agree well with the first-order theory given in part I.

When the disturbance has attained a considerable amplitude the first-order equations cease to apply and it changes from a wave into a form which has the appearance of large round-ended columns of air extending into the liquid and separated by narrow sheets of liquid. The air columns attain a steady velocity relative to the accelerating liquid and continue to penetrate into the liquid until the lower surface of the liquid is reached.

In spite of these very large surface disturbances, the main body of liquid below them is accelerated as though they did not exist.

1. INTRODUCTION

The experiments described in this paper were designed to test the validity of Sir Geoffrey Taylor's first-order theory (given in part I, Taylor 1950) for the instability of an initially disturbed sheet of liquid when subsequently accelerated. It is shown in part I that the increase in amplitude of a two-dimensional disturbance accelerated downwards at an acceleration g_1 is given by

$$\frac{\eta}{\eta_0} = \cosh \sqrt{\left\{ \frac{4\pi s}{\lambda} \frac{(g_1 - g)}{g_1} \frac{(\rho_2 - \rho_1)}{(\rho_1 + \rho_2)} \right\}}, \quad (1)$$

where η , η_0 , s , λ , ρ_1 , ρ_2 , g_1 and g are as defined in part I.

The experimental apparatus was so constructed that a range of accelerations of from $3.0g$ to $140g$ (g being 32.2 ft./sec.²) could be applied to various depths of liquid of from $\frac{3}{8}$ to 20 in.

2. DESCRIPTION OF APPARATUS AND ITS OPERATION

A photograph of the acceleration apparatus is given as figure 1. The liquid is initially supported in a vertical channel whose cross-section is $2\frac{1}{2} \times \frac{1}{2}$ in. by means of a thin brittle diaphragm inserted between the flanges A joining together the two glass-sided observations ducts B and C . After many trials with glass, waxes, paper, etc., a technique was developed for making diaphragms which would break at a pressure equivalent to as little as 2 in. of water and yet would not leak under a steady pressure equivalent to 1 in. of water. These diaphragms were made of shellac of thickness 0.001 to 0.007 in.

The portion of the apparatus above the liquid is airtight and that below the liquid is also rendered airtight when a glass disk is clamped between the flanges *D* at the bottom of the apparatus. Air is then admitted through the pipes *E* and *F* from a single source of compressed air, so that the pressure of the air above the liquid is raised at the same rate as the pressure of the air below the liquid.

When the desired pressure has been reached, the volumes of air above and below the liquid are isolated from each other and from the compressed air supply. At the

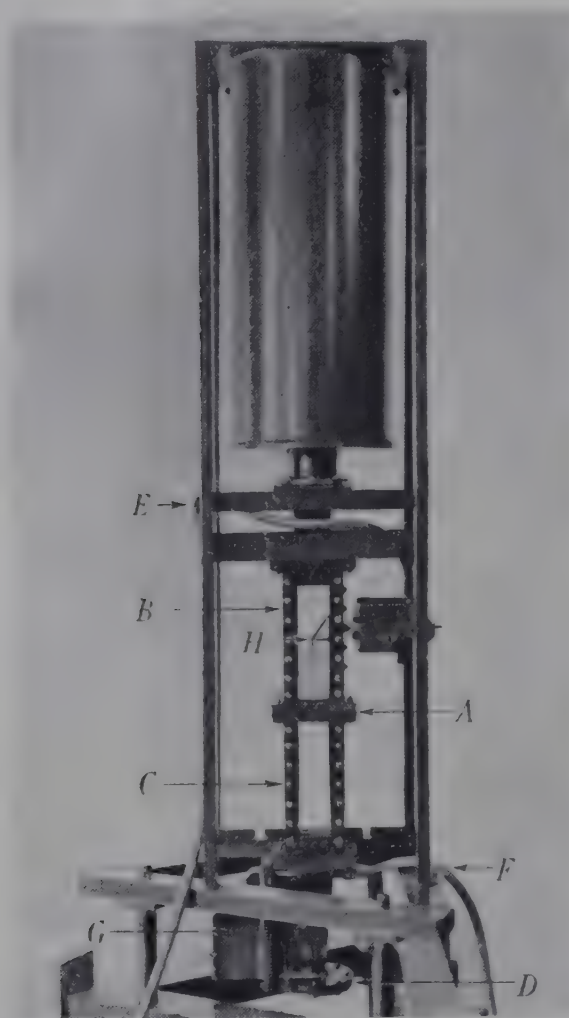


FIGURE 1. The acceleration apparatus.

appropriate instant, a mechanism inside the cover *G* is released electrically which breaks the glass disk held between the flanges *D*. The resulting rarefaction wave travels up the apparatus, and the shellac diaphragm breaks completely and allows the liquid to accelerate downwards under the action of the compressed air above it.

By this technique, a column of liquid of height 9 in. can be accelerated at accelerations of up to $75g$ and smaller columns of liquid at proportionately higher accelerations. The top surface of the liquid can be given an initial surface disturbance by means of the agitating paddle *H* which is actuated by a motor and mechanism in the box seen to the right of the observation duct *B*.

The downward motion of the liquid is recorded by means of shadow photography through the two pairs of toughened glass windows *B* and *C*. The arrangement of the

photographic apparatus is shown in figure 2. The illumination is provided by discharges of charged condensers across four spark gaps which are arranged in two groups. Two spark gaps separated by a thin sheet of mica are placed on the optical axis of a condensing lens J which concentrates their light flashes into the lens of a quarter-plate camera K fitted with a $f/8$ lens of 6 in. focal length. The effective area of the condensing lens is increased by the use of a spherical mirror L which produces an image of the spark gaps a few inches in front of their actual position. By this means a 7 in. length of the upper window B is illuminated.

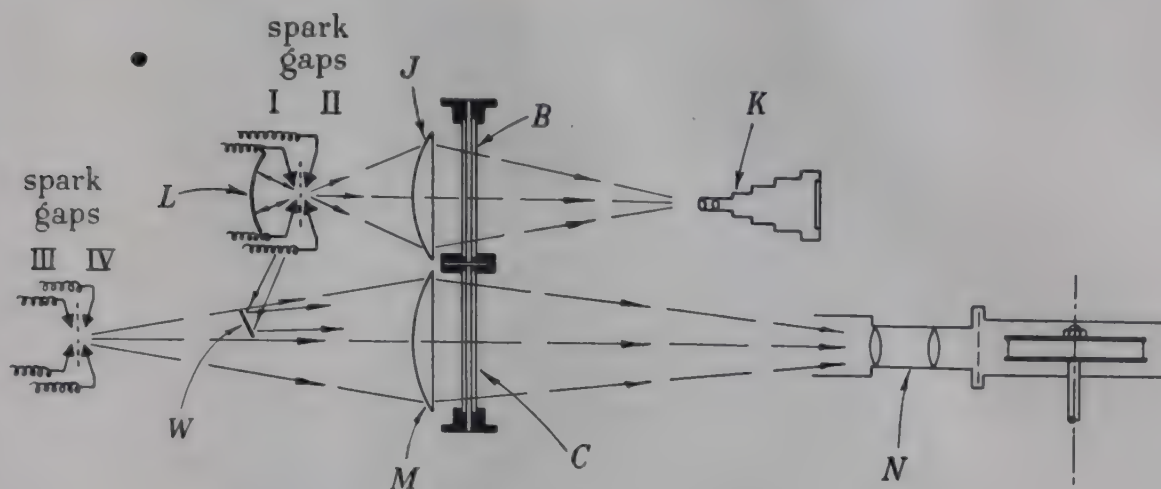


FIGURE 2. The optical system.

The lower window C is illuminated by another pair of spark gaps placed on the optical axis of a condensing lens M which concentrates the beam to pass through the $f/2.5$ 7 in. lens N of a rotating drum camera. This lens M covers the whole 9 in. length of the window C .

A photograph of the drum camera is reproduced as figure 3. It consists of an aluminium drum O which carries a length of unperforated 35 mm. film clamped on the outside and can be rotated at speeds up to 7000 r.p.m. by means of the motor P . This motor also drives a shaft carrying a balanced arm Q whose tip passes close to a series of contact screws R mounted on a ring of insulating material. These contact screws are so arranged circumferentially that a set of twenty-four photographs can be spaced evenly around the periphery of the drum O and also that twelve additional photographs can be taken in between the twenty-four previously exposed on the same length of film. The motor P also drives the generator of an electrical speed indicator which gives continuous indication of the motor speed.

The electrical circuits employed to produce the illuminating sparks are shown in figure 4. A valve rectifier is used to provide a source of 5500 V d.c. which is used to charge up a number of separate condensers C_1 and C_3 through high resistances.

The sequence of events making up an experiment are controlled by the pendulum S which is mounted on horizontal pivots and initially in a horizontal position. As it swings downwards it operates a series of switches (not all shown on figure 4) mounted on an arc of insulating material.

Its first action is to make a circuit which energizes two solenoids which open the camera shutters. It then electrically actuates the mechanism to break the glass disk at the bottom of the apparatus approximately 30 msec. later.

Just before the liquid starts to move, the pendulum touches the contact screw *T* which causes one condenser to discharge across spark gap I which records the initial state of the top surface of the liquid in the plate camera. Soon afterwards the pendulum makes contact with a brass arc *U* whose length is such that the time of

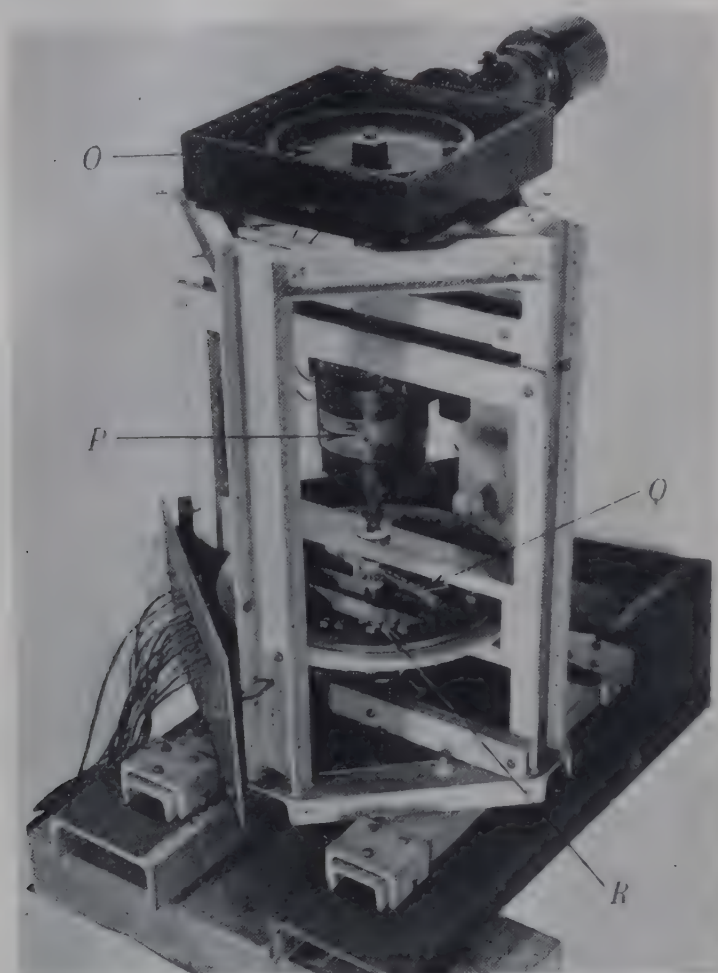


FIGURE 3. The drum camera.

contact is equivalent to one revolution of the drum camera. Whilst contact is made here, twelve condensers discharge across spark gaps II and III in turn when the rotating arm brushes past the appropriate contact screw *R*. The six sparks across spark gap II record profiles of the downward motion of the top liquid surface on the plate in the plate camera and the other six across spark gap III record the motion of the bottom liquid surface on the length of film.

The succession of up to seven profiles recorded on the stationary plate can be separated easily because the liquid leaves a film on the surface of the windows after the passage of the liquid which interrupts the light beam. The plate on development shows a series of areas of differing densities, the boundaries of which correspond to the top surface profiles.

After this set of photographs, more are taken on the drum camera when the pendulum touches the arc V which allows the rotating arm Q to time twenty-four discharges across spark gap IV. Subsequently, the pendulum causes the camera shutters to be closed and is prevented from swinging back.

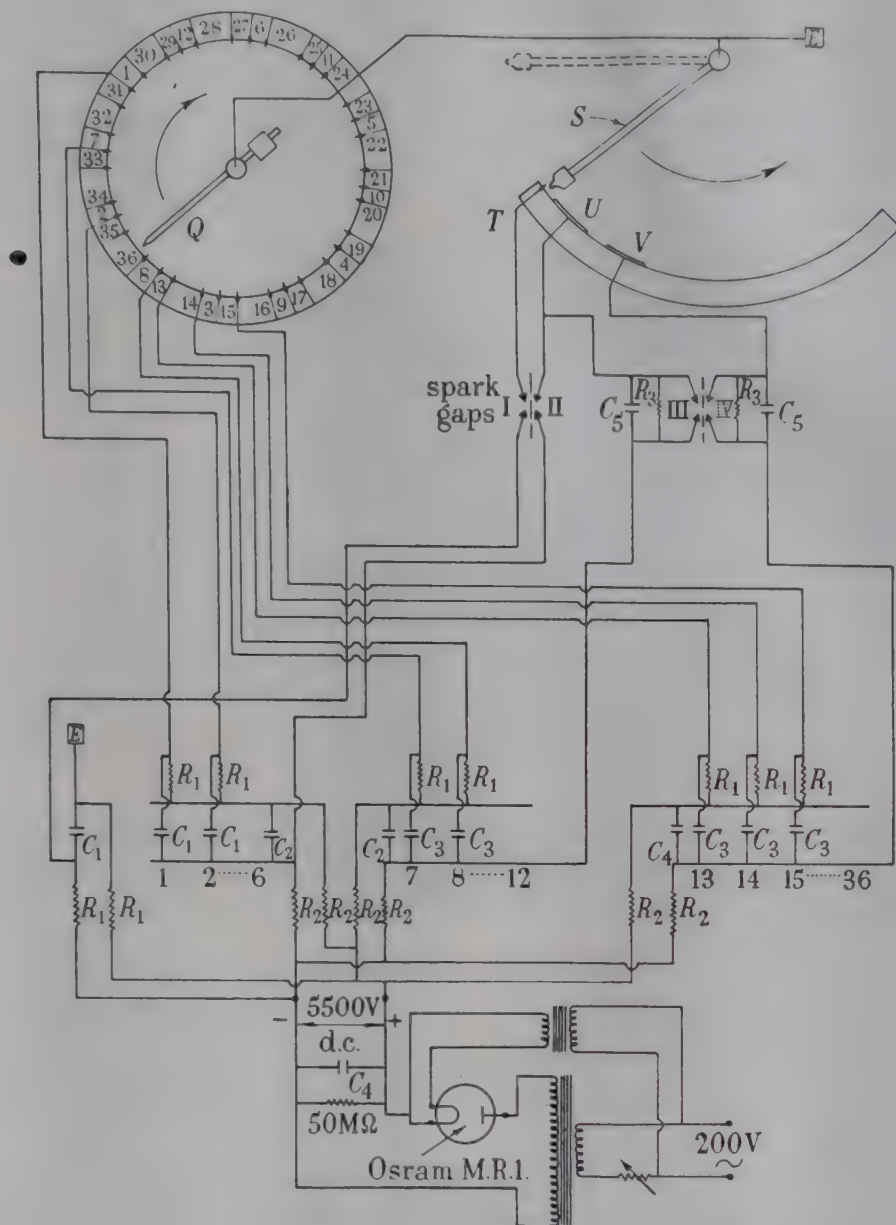


FIGURE 4. The electrical circuits.

$R_1 = 10 \text{ M}\Omega$, $R_2 = 10 \text{ k}\Omega$, $R_3 = 1 \text{ k}\Omega$, $C_1 = 0.1 \mu\text{F}$, $C_2 = 0.5 \mu\text{F}$, $C_3 = C_5 = 0.2 \mu\text{F}$, $C_4 = 1 \mu\text{F}$.

The duration of the discharges across spark gaps III and IV is kept as low as possible by the provision of the condensers C_5 which are situated a few inches from the spark gaps. The spark gaps are 0.030 in. long between pointed copper electrodes.

In order to establish the instant at which each exposure is made in the plate camera, a small beam of light is reflected from spark gaps I and II by the plane mirror W (see figure 2) to illuminate part of the lower pair of windows C .

The plates used in the plate camera were Ilford Ordinary in the majority of experiments, but recourse had to be made to Kodak 0800 plates in a few experiments. The film used in the drum camera was Ilford Recording Film 5G 91.

3. THE BEHAVIOUR OF LIQUID SURFACES AS RECORDED EXPERIMENTALLY

A considerable number of experiments were made using water as the liquid accelerated under a variety of conditions. Of these, five have been selected for discussion here as being representative of the experiments carried out with water.

The first one is film 2 (reproduced as figure 5, plate 5), which was taken before the photographic arrangement had been fully developed. Only one spark gap was in use which recorded shadow pictures on the drum camera of that portion of the acceleration apparatus extending from $1\frac{1}{2}$ in. above the bottom of the upper windows *B* to $2\frac{1}{2}$ in. above the bottom of the lower windows *C*.

In all the experimental records reproduced, the profiles of the upper water surface have been numbered in order of time with the initial top surface termed 0. The time that had elapsed since the start of the acceleration at the instant each profile was recorded is also given.

The acceleration depicted in film 2 was equal to $20.7g$, and this was applied to 1.14 in. of water whose top surface had been given an approximately sinusoidal disturbance of initial amplitude 0.034 in. and initial wave-length 0.83 in. Profiles 1 to 3 show a rapid increase in amplitude of the top surface disturbance. The top of the white portion of the photographs showing the top-surface profile is the envelope of the lowest points to which the compressed air had descended at that instant, and small quantities of water are present above this profile in contact with the sides of the channel.

The bottom surface of the water is first recorded in profile 4 as a black line. This indicates that the breakage of the shallac diaphragm produces very little disturbance of the bottom surface.

The subsequent profiles 5 to 8 show the top surface having the form of a number of round-ended columns of air separated by narrow upstanding columns of water. The top surface undergoes comparatively little change of profile during this period, whilst the air column on the left descends relative to the water until it finally reaches the bottom surface.

The second one is film 7, which is reproduced as figure 6, plate 6, in the form of an enlargement of the superimposed profiles recorded on the plate and a series of enlargements of the images recorded on the drum camera. Films 28 and 23 are similarly reproduced as figures 7 and 8, plates 7 and 8.

The acceleration depicted in film 7 was equal to $52.3g$ acting on 5.29 in. of water, and the initial disturbance was approximately sinusoidal of initial amplitude 0.044 in. and initial wave-length 1.49 in. The increase in amplitude shown in the profiles 1 to 5 is very large. These profiles indicate the envelope of the lowest points to which the compressed air has descended as in the case of film 2, figure 5.

The enlargements from the length of film show considerably less detail in the profile owing to differences in image size and emulsion granularity between the two records. They show that the progressive change from a number of surface troughs and peaks to a smaller number continues until one round-ended column of air remains symmetrically disposed in the acceleration channel.

The characteristics of film 7 are also shown by film 28, figure 7, where 6.80 in. is shown as it was accelerated at 42.8g after being given an initial disturbance of amplitude 0.025 in. and wave-length 0.40 in. The subsequent profiles are in all respects very similar to those shown in film 7 (compare profiles 5 of each film).

Film 23, figure 8, shows what happens when the top liquid surface is accelerated without first applying an initial disturbance. The slight surface tension curvature at the sides of the channel and other unavoidable sources of small disturbances are sufficient to start the instability, but no appreciable top-surface amplitude appears until the water has descended a couple of inches. The later pictures taken on the

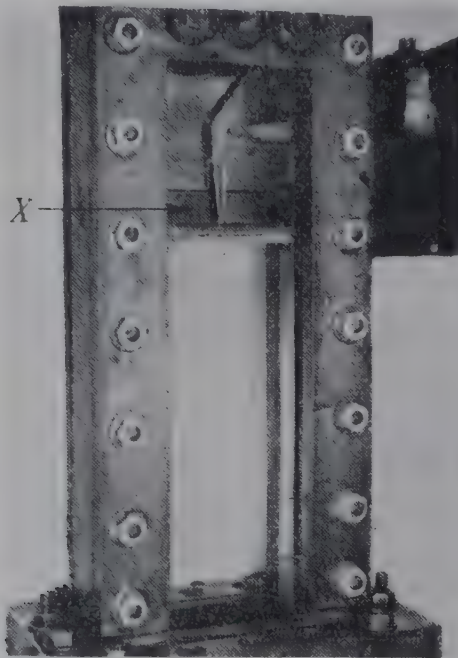


FIGURE 9. Arrangement for accelerating small depths of water.

drum camera show a considerable amplitude similar in all respects to earlier stages of films 7 and 28.

The fourth film reproduced as figure 10, plate 9, was taken using a slightly different technique. It was desired to observe the behaviour of a small depth of water, and to this end a small tank having two thin mica sides and a shellac diaphragm forming its base was inserted between the windows *B* (see figure 1). The arrangement is shown in figure 9, in which the tank *X* can be seen containing a small depth of water with the agitating paddle in position to provide the initial surface disturbance. The acceleration apparatus as a whole was lowered so that the windows *B* were in position for illumination by spark gaps III and IV and recording on the drum camera.

The film 37, figure 10, taken by this method shows the inherent stability of the bottom surface of the water very well. Until the depth of the water (initially 0.62 in.) has been reduced by the steady penetration of the air columns through the water to about $\frac{1}{4}$ in. (profile 4), the bottom surface has remained quite flat. Thereafter it is depressed downwards in front of each air column until the water separating the two volumes of air has been reduced to a thin film. The film of water is then blown out

ahead of the bottom surface (causing the dark areas below it in profile 5) and finally bursts (profile 6).

The behaviour of the top surface shown in film 37 is similar to that observed in films 7 and 28.

The instability as observed experimentally is not two-dimensional owing to the presence of the channel sides and cannot immediately be compared to the theoretical two-dimensional discussion of part I until the magnitude of the deviation from the two-dimension case has been assessed. In order to do this, a small number of experiments were made in which the unstable surface was observed in two vertical planes at right angles.

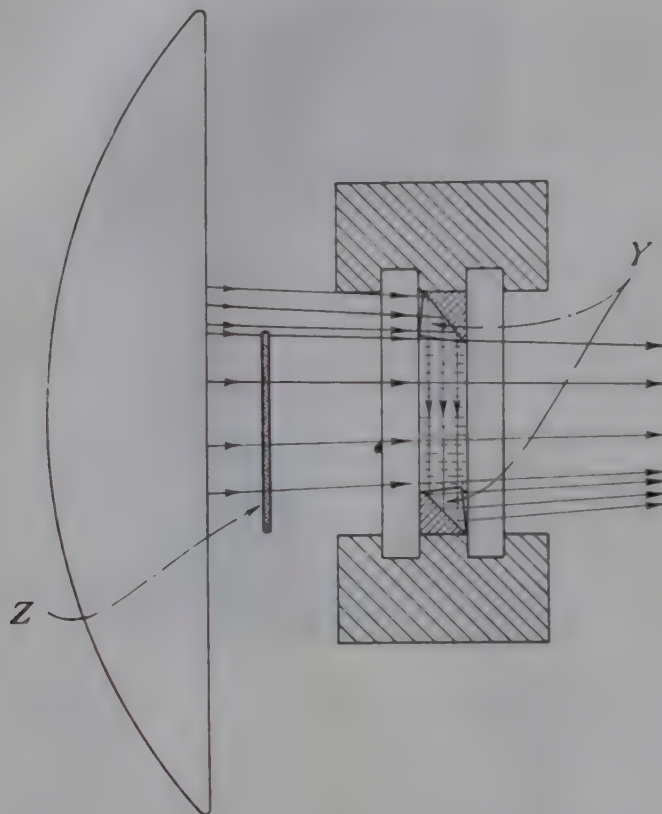


FIGURE 11. Arrangement of cross-viewing prisms.

This was achieved by the use of two prisms inserted in the acceleration channel as shown in the horizontal section reproduced as figure 11. The two prisms *Y* divert part of the illuminating beam of light to traverse the channel parallel to the faces of the glass windows and then to rejoin the direct beam on its way to the plate camera. A neutral density filter *Z* was interposed in the direct beam to make the two beams of light of comparable intensity.

The records obtained by the use of this arrangement show that when the top surface has descended about 4 or 5 in. (and has a direct profile of large radius—of the order of 0.5 to 1.0 in.) the cross-profile shown by the prisms had a radius of curvature of 0.225 in. As the half-width of the channel is only 0.250 in., the air column is very nearly in contact with the glass windows and therefore the instability can be regarded as mainly two-dimensional up to this stage.

A few experiments were made to observe the instability of an interface between two non-miscible liquids. One such film is given as figure 12, plate 10, in which

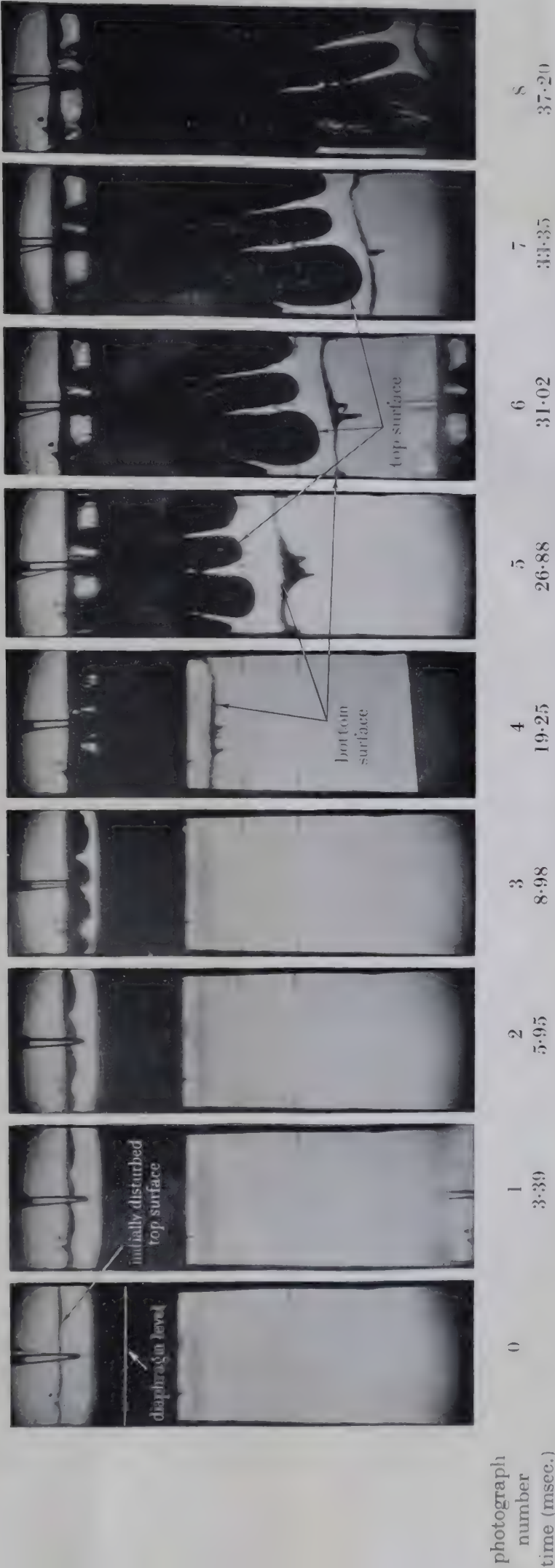


FIGURE 5. Acceleration of water; film 2. Depth of water 1.14 in.; initial acceleration = 20.7g.

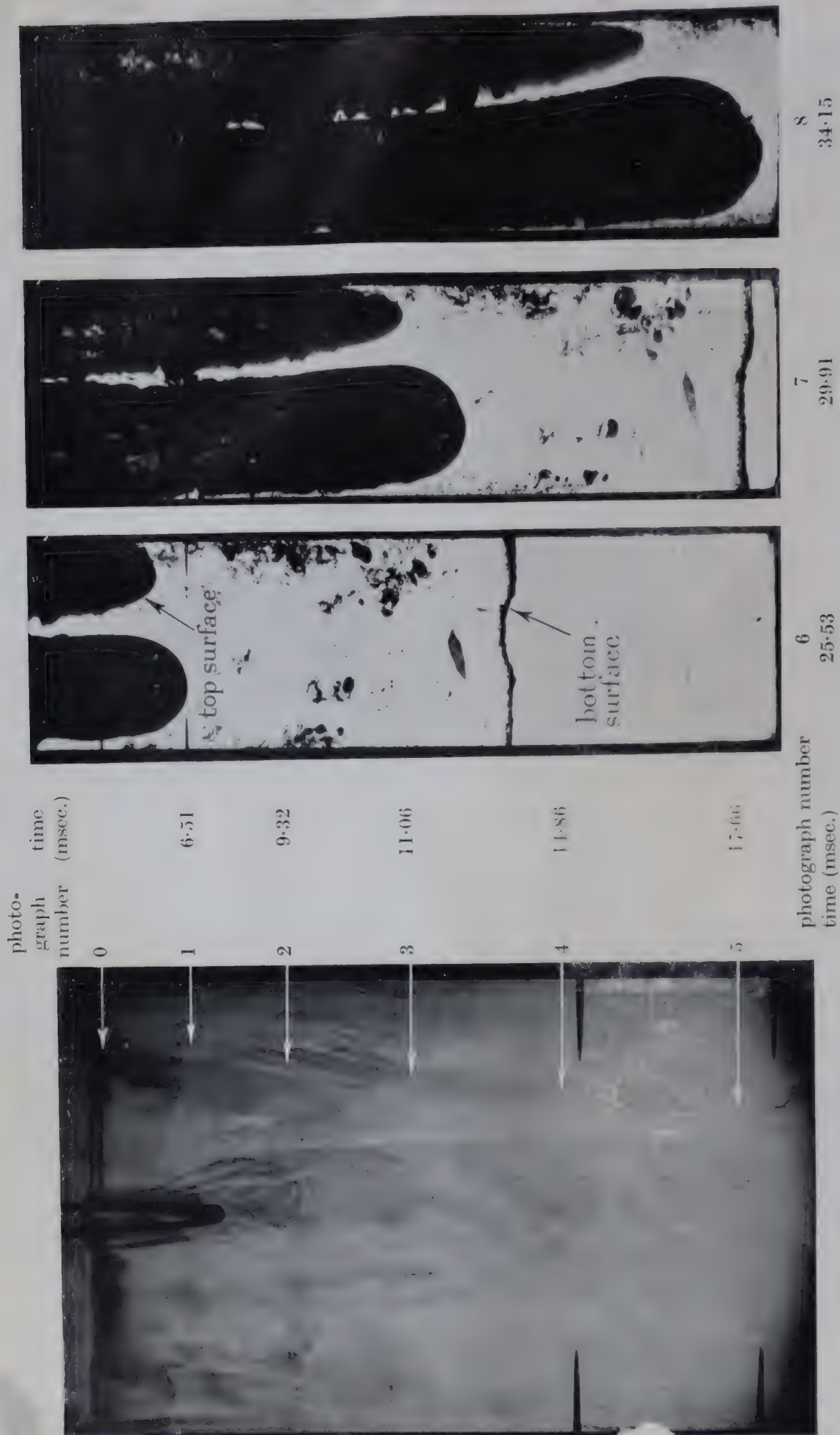


FIGURE 6. Acceleration of water; film 7. Depth of water 5.29 in.; initial acceleration 52.3g.

photo-
graph
number

time
(msec.)



0

1

2

3

4

5

3.04

10.61

13.89

17.13

20.39

photograph number
time (msec)

6

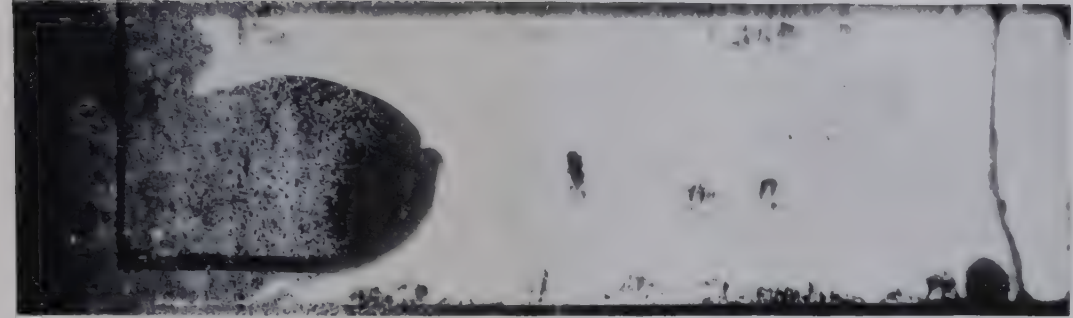
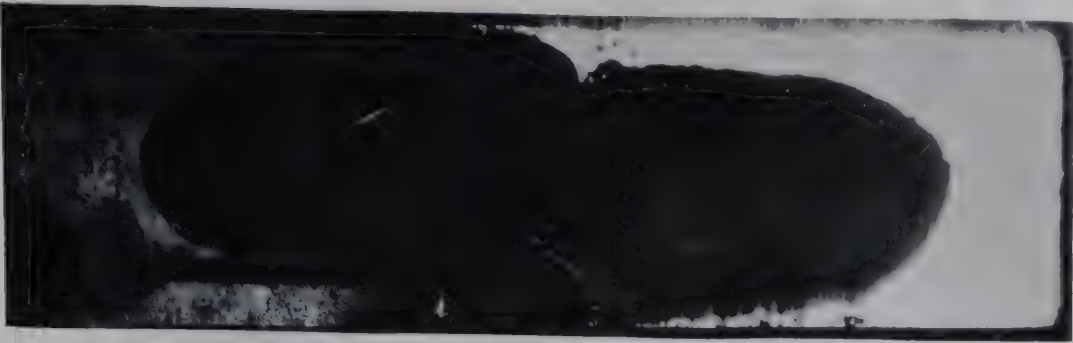
33.15

7

36.41

8

39.69



Acceleration of water; film 28. Depth of water = 6.80 in.; initial acceleration 42.8g.

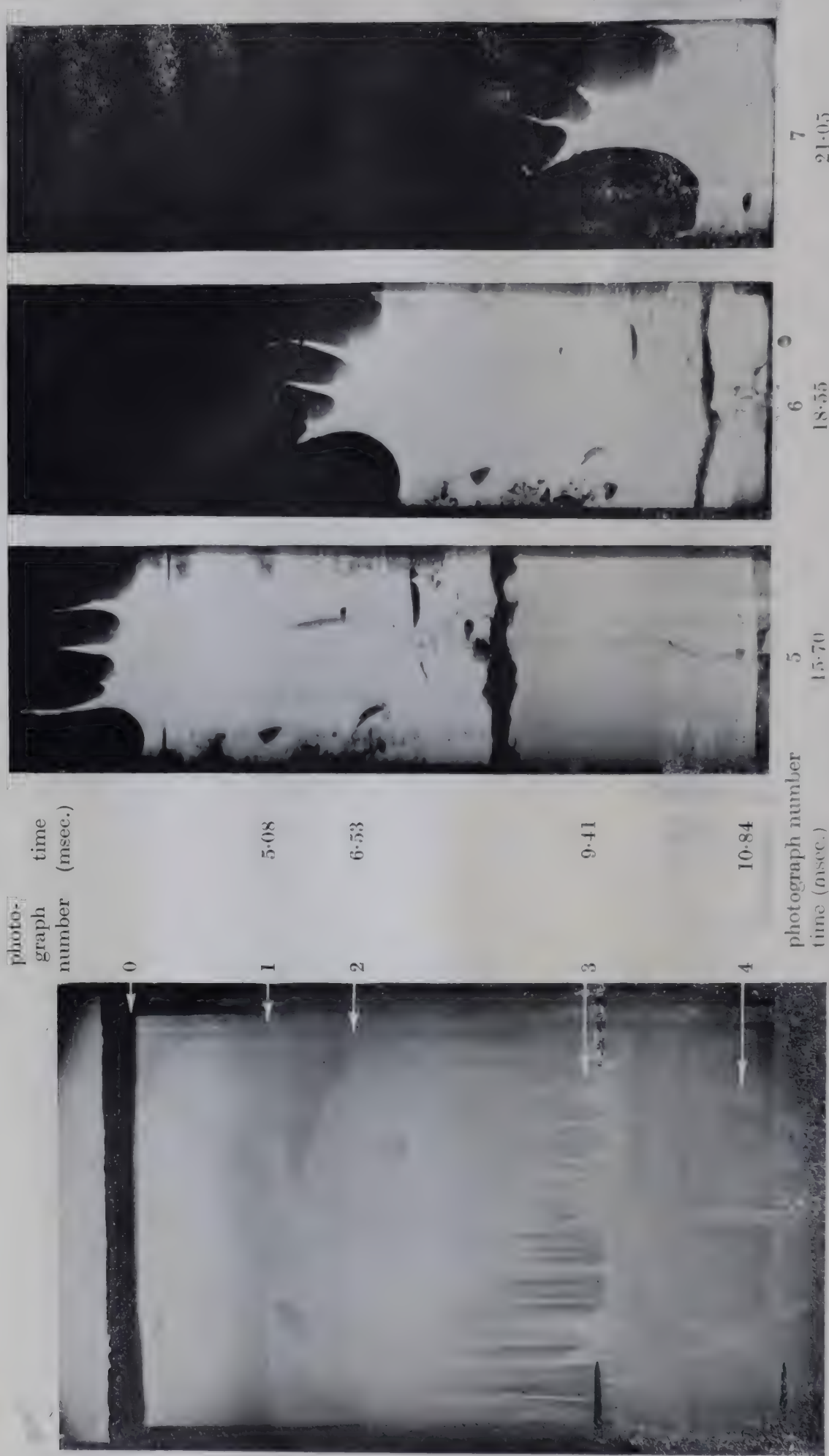


FIGURE 8. Acceleration of water with no applied initial disturbance; film 23. Depth of water = 5.20 in.; initial acceleration = 131.9g.

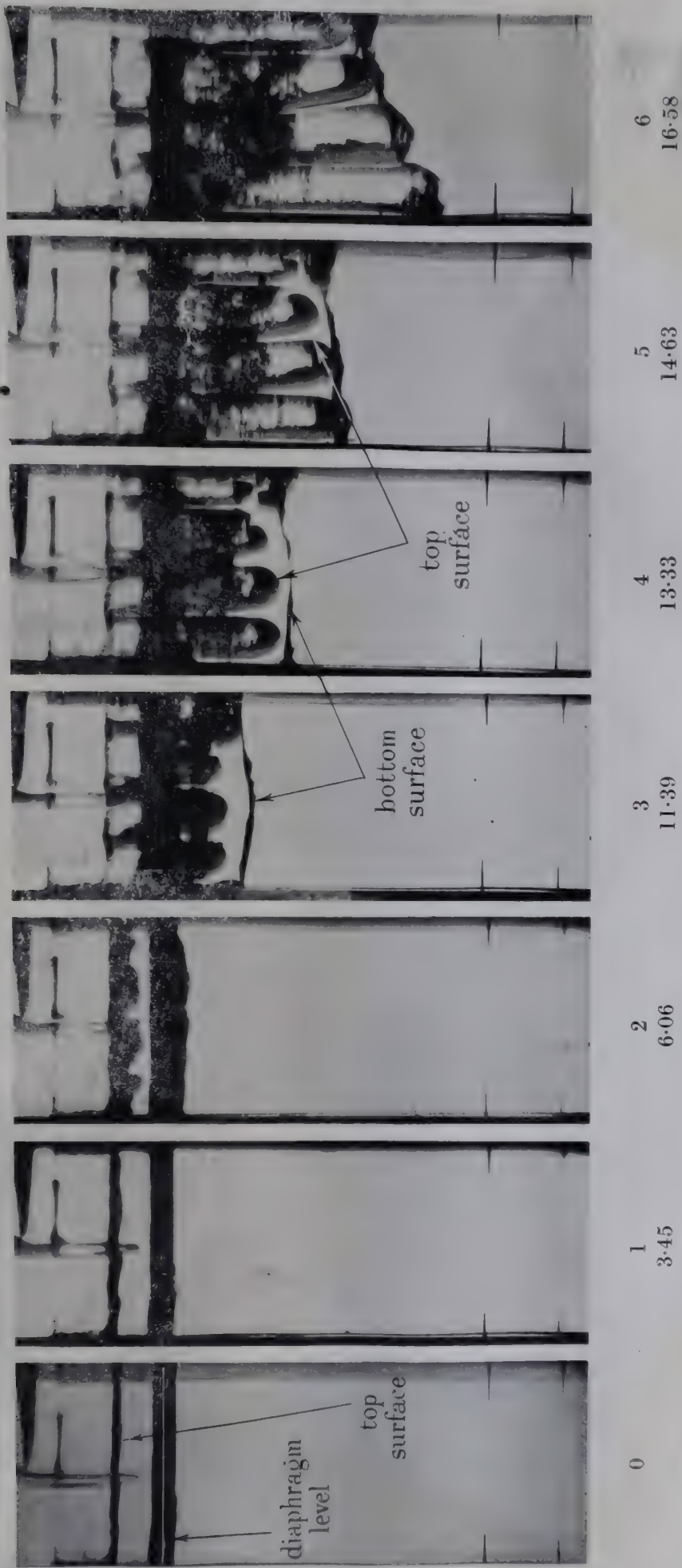
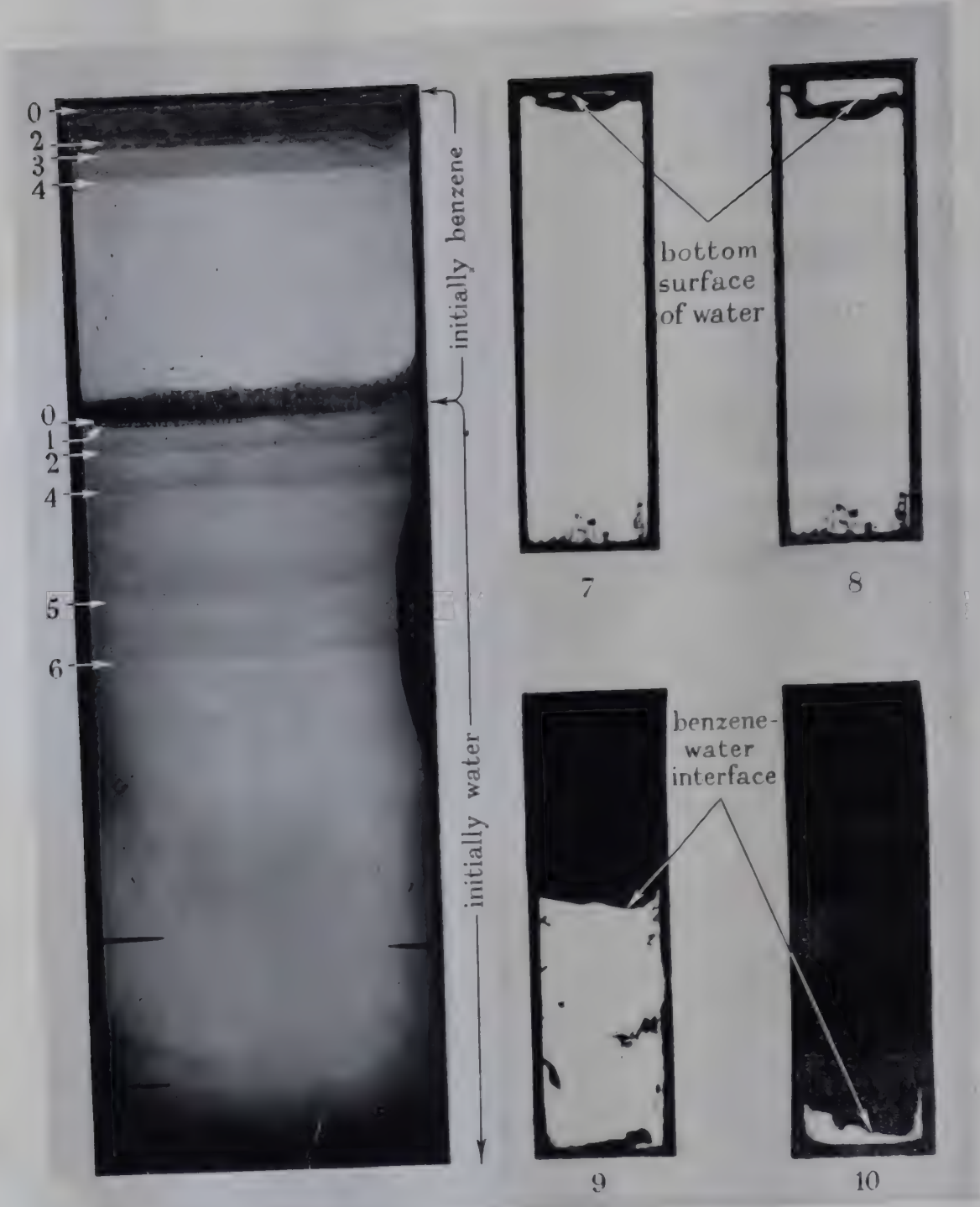


FIGURE 10. Acceleration of a small depth of water; film 37. Depth of water = 0.62 in.; initial acceleration = $35.8g$.

photograph number
time (msec.)

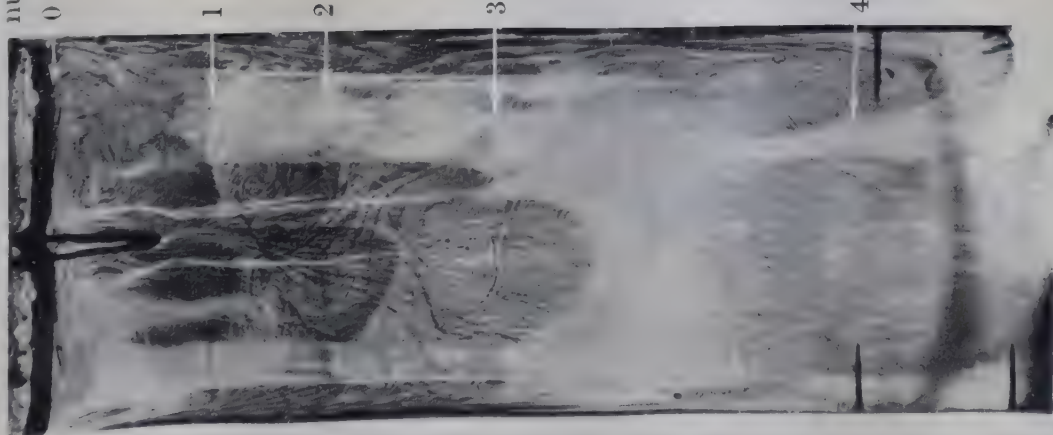


ph number	1	2	3	4	5	6	7	8	9	10
(sec.)	5.10	7.57	8.30	10.14	15.21	16.18	12.58	14.62	42.74	49.86

2. Acceleration of benzene-water interface; film 45. Depth of benzene = 2.10 in.
depth of water = 4.70 in.; initial acceleration = 32.1g.

photo-
graph
number

time
(msec.)



0

1

2

3

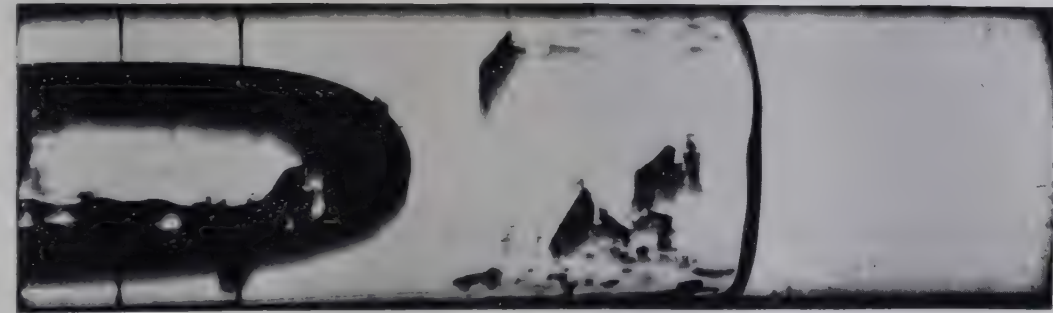
4

9.71

12.86

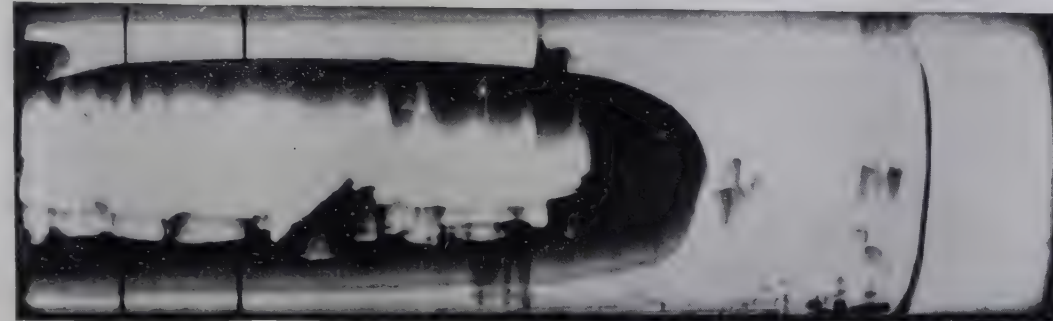
16.04

22.35



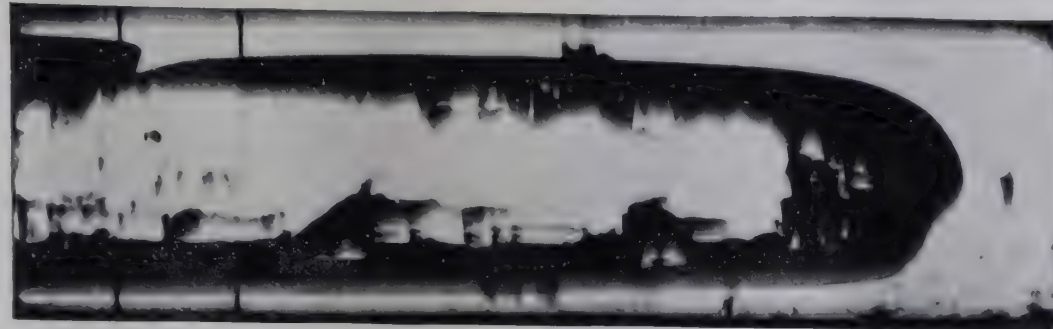
5

28.30



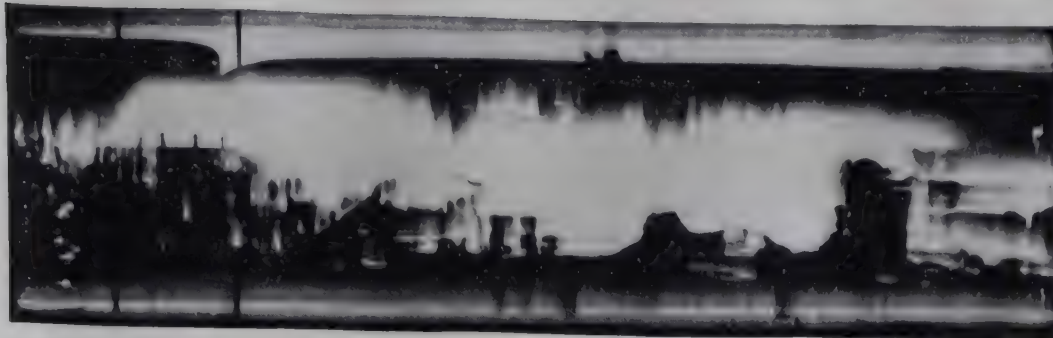
6

31.46



7

33.05

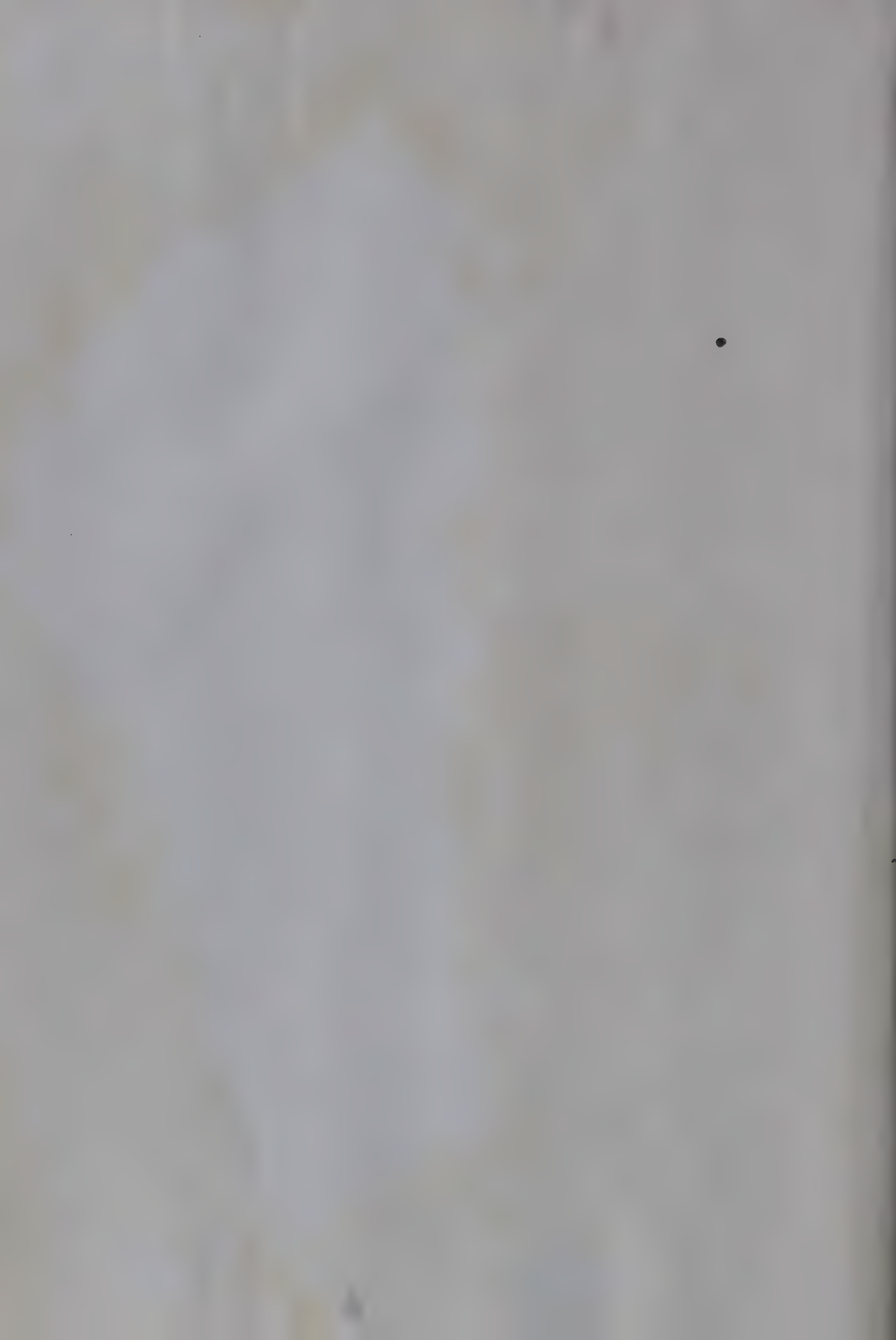


8

34.65

graph number
(sec.)

ion of glycerine; film 67. Depth of glycerine = 7.88 in.; initial acceleration = 41.4g.



2.10 in. of benzene resting on top of 4.70 in. of water were accelerated at $32.1g$. The top surface of the benzene was stabilized by floating a thin strip of balsa wood on it, and an initial slight disturbance given to the interface by means of a pulsating diaphragm placed in the wall of the channel.

The early profiles of the interface (1 to 6) show negligible change, but the two enlargements from the drum camera film (profiles 9 and 10) show that an appreciable amplitude has developed. This decrease in the rate of development of the instability is in full agreement with equation (1), representing the conclusions of the theoretical discussion given in part I.

Finally, a number of experiments were made with glycerine. A representative film is reproduced as figure 13, plate 11. The instability develops in the early stages very similarly to that observed with water, but the later stages differ in that a large quantity of glycerine is left adhering to the sides of the channel, with the result that the air penetrates through the liquid at a much faster rate. Over half of the quantity of glycerine introduced into the apparatus proceeds to drain out of the apparatus subsequent to the actual acceleration and the 10 sec. immediately after.

4. ANALYSIS OF THE EXPERIMENTAL RECORDS

The time intervals between successive photographs were established by measuring the distance separating the images or timing marks concerned on the length of film from the drum camera and by reference to the speed of rotation of the drum at the time of the experiment.

The negatives were examined by means of an optical enlarging system, and the mean horizontal levels of the various profiles of the top and bottom surfaces determined. The distances of these horizontal levels from the initial positions of the surfaces were taken as the distances that the top and bottom of the liquid had descended (s_t and s_b respectively).

The acceleration of the liquid was established by plotting $\sqrt{s_t}$ and $\sqrt{s_b}$ against the time intervals. This gave the time of starting of the acceleration from which the values of time t were then made to refer. The plottings of these points are reproduced for films 7 and 23 as figures 14 and 15 respectively. It will be seen in both cases that the acceleration of the bottom surface of the water is remarkably uniform. The top surface accelerates at about 10 to 15 % faster than the bottom surface during the acceleration of water. This is caused by the quantity of water that is left behind on the sides of the channel reducing the effective depth of the main body of water. This effect is accentuated in the case of small depths of water, and is much greater (of the order of 45 %) in the case of glycerine.

With each experiment, the amplitude of the most prominent trough in the disturbance shown in each profile has been measured; and the measurements have been plotted as $\cosh^{-1} \frac{\eta}{\eta_0}$ against $\sqrt{\left\{ \frac{s(g_1 - g)(\rho_2 - \rho_1)}{\lambda g_1(\rho_1 + \rho_2)} \right\}}$, using distinctive symbols for the points derived from the various types of experiments. The result is reproduced as figure 16, on which the straight line representing equation (1) is also shown. It is

seen that the points having values of $\sqrt{\left\{ \frac{s(g_1 - g)(\rho_2 - \rho_1)}{\lambda g_1(\rho_1 + \rho_2)} \right\}}$ less than about 0.8 are in excellent agreement with equation (1), and that each class of experiment, whether

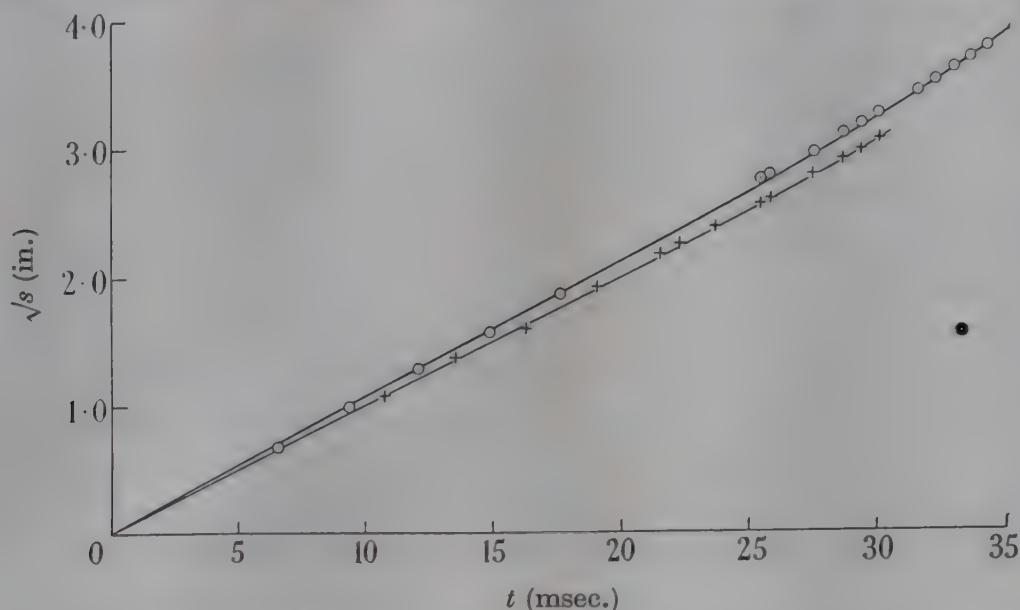


FIGURE 14. $\sqrt{s}$ against t for film 7: $\odot$, derived from top surface; $+$, derived from bottom surface.

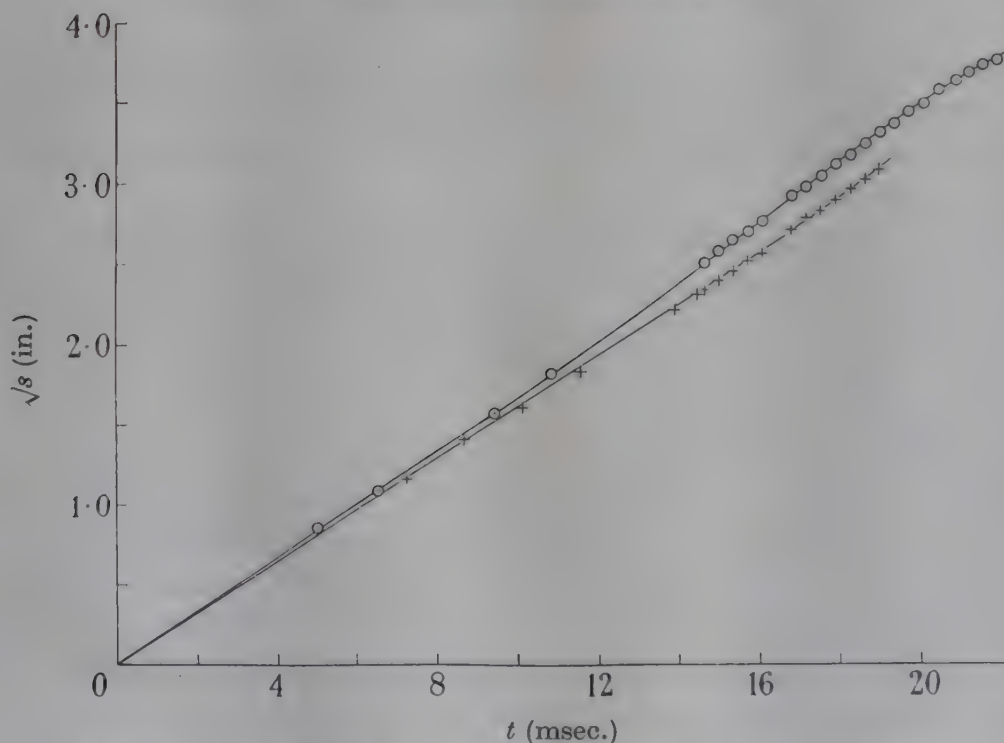


FIGURE 15. $\sqrt{s}$ against t for film 23: $\odot$, derived from top surface; $+$, derived from bottom surface.

dealing with water, glycerine or two liquids, gives rise to points which are scattered throughout the band of experimental points.

To enable further conclusions to be drawn, the points derived from each experiment were plotted and a smooth curve drawn through them. These curves have been combined into two diagrams, figure 17 showing all those derived from liquids of low viscosity and figure 18 showing those derived from the experiments with glycerine.

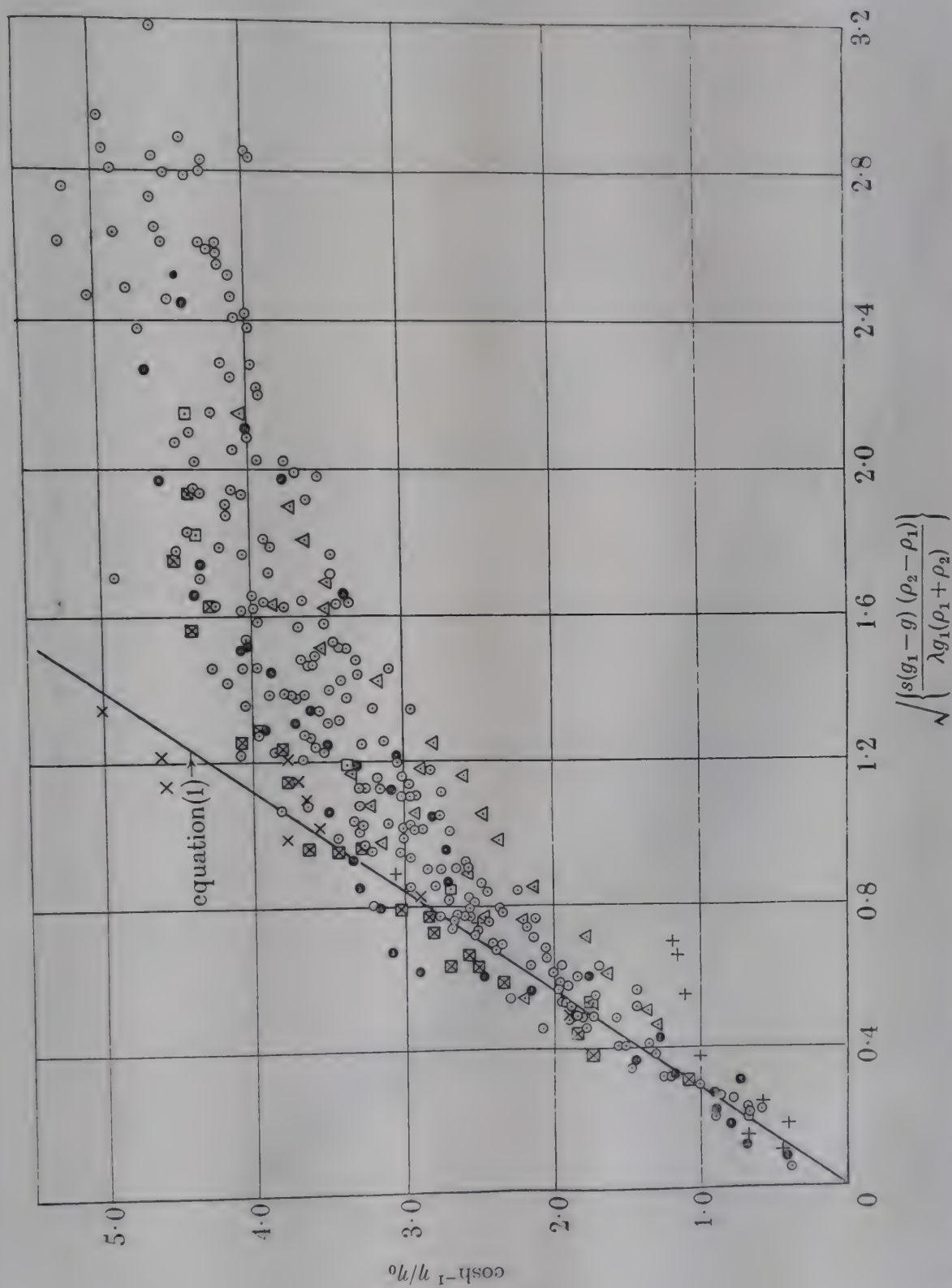


FIGURE 16. Analysis of unstable surfaces—experimental points: $\odot$, $g_1 > 10g$ and $h > 1.5$ in.; $\triangle$, $g_1 < 10g$ and $h > 1.5$ in.; $\square$, air-benzene surfaces; $\times$, water-carbon tetrachloride surfaces; $+$, benzene-water surfaces; $\times$, air-glycerine surfaces.

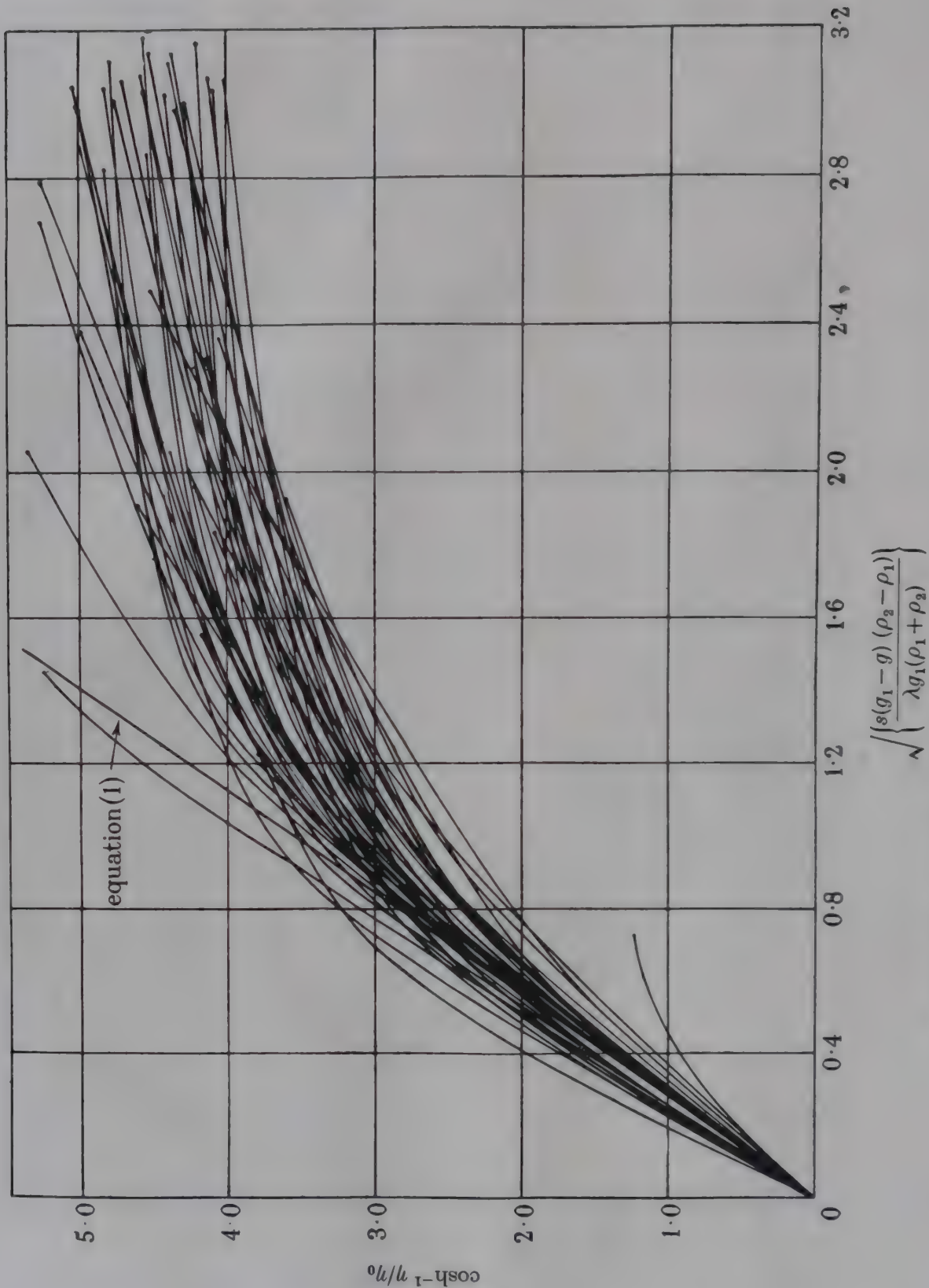


FIGURE 17. Analysis of unstable non-viscous surfaces—experimental curves.

The initial slopes of the curves shown in figure 17 show a considerable variation, but are equally distributed on either side of the line corresponding to equation (1).

The experimental points through which these curves were drawn depend for their position vertically on the value assigned to η_0 . This quantity was the smallest one which had to be measured (varying from 0.005 to 0.336 in.), and unavoidable errors made in this measurement will give rise to the variation of the initial slopes of the curves of figure 17.

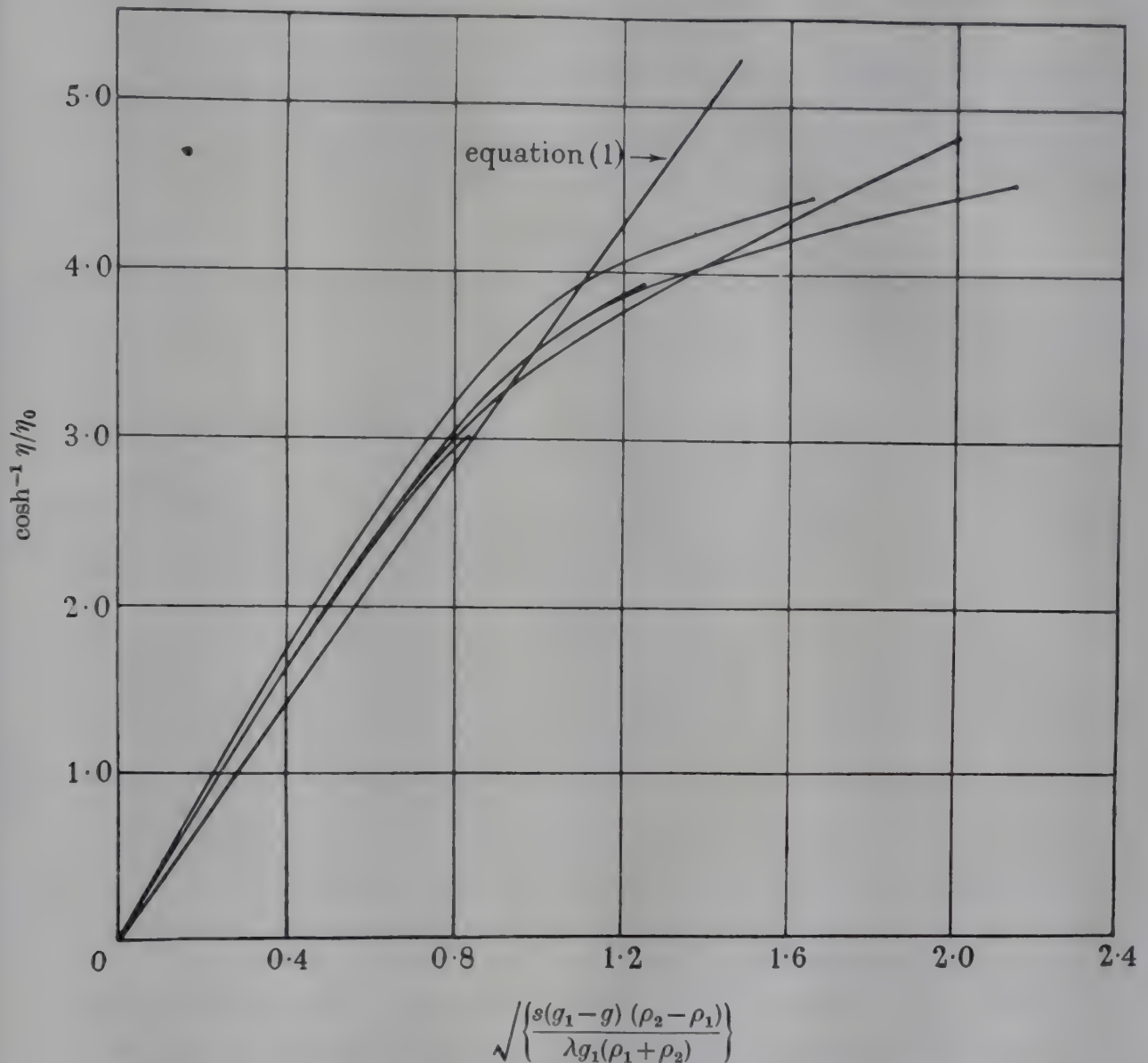


FIGURE 18. Analysis of unstable air-glycerine surfaces—experimental curves.

It therefore appears that the initial rate of development of the instability of the surfaces of liquids of low viscosity is correctly given by equation (1). From figure 18, it is concluded that the instability of air-glycerine surfaces develops slightly faster than that given by equation (1), but this is probably due to the viscous drag exerted by the channel sides and would be much reduced in the case of a channel of larger cross-section.

The point on each curve at which the slope is equal to half that of the line representing equation (1) has been determined, and the ratio of η_{c_1} (the amplitude at this point) to λ has been calculated. This ratio η_{c_1}/λ represents the shape of the unstable

disturbance at the time when it is ceasing to increase at the rate given by the first-order theory.

The values of η_c/λ range from 0.11 to 0.96 with 75 % of them between 0.23 and 0.63. As these values will be dependent on the positions of the experimental points of figure 16 (which are determined by the estimation of the value of η_0 as explained above), it seems that these results may be interpreted roughly as showing that the instability follows the first-order theory until it has attained an amplitude of about 0.4λ .

The succeeding stages of the instability consist of round-ended columns of air penetrating steadily through the liquid with little change of profile. It has been shown by Taylor & Davies (1950) that air will ascend up a vertical cylinder of radius a initially full of stationary liquid at a constant velocity which is given by

$$v = 0.48 \sqrt{ga}. \quad (2)$$

Similarly, it would be expected that the velocity appropriate to the two-dimensional case would be given by an expression of the form

$$v = C_1 \sqrt{gr_1}, \quad (3)$$

where r_1 is the radius of curvature of the tip of the air column and C_1 is a constant, not necessarily the same as that in equation (2).

The air columns penetrating into the accelerating liquid will in effect be moving relative to the liquid in a potential field of value $(g_1 - g)$. The relative velocity v will therefore be given by

$$v = C_1 \sqrt{(g_1 - g)r_1}. \quad (4)$$

The experimental records were again examined when the minimum distance between the bottom surface and the troughs of the top surface was determined in each photograph. These minimum distances were then plotted against time and a specimen graph is reproduced as figure 19. It is seen that the rate of decrease of depth of the liquid is quite low, until at about $9\frac{1}{2}$ msec. after the start of the acceleration it abruptly changes to a much higher and constant value. This constant velocity v has been determined in this manner for each suitable experimental record, and the mean radius of the tip of air column used to determine v over the time during which v is constant has also been measured as $\bar{r}_1$.

From these measurements of the experiments with water v has been plotted against $\sqrt{(g_1 - g)r_1}$ to give figure 20. The straight line thereon represents the best value of C_1 in equation (4) which is found to equal 1.11. The considerable scatter of the points is partly due to errors arising in the measurement of r_1 and partly due to the assumption that the experiments are truly two-dimensional. It has been shown that the velocity is accurately proportional to $\sqrt{(g_1 - g)}$ but not to $\sqrt{r_1}$. Equation (4) with $C_1 = 1.11$ over-estimates the value of v when r_1 is small (about 0.25 in.) and under-estimates the value of v when r_1 is large (about 0.9 in.).

Similar measurements of air-glycerine surfaces give very much higher values of v , which require the value of C_1 to be between 3.5 and 5.5 (according to the value of the viscosity) for equation (4) to predict the measured values of v . This increase in the penetration velocity is due to the influence of the sides of the channel and would not be observed in channels of larger section.

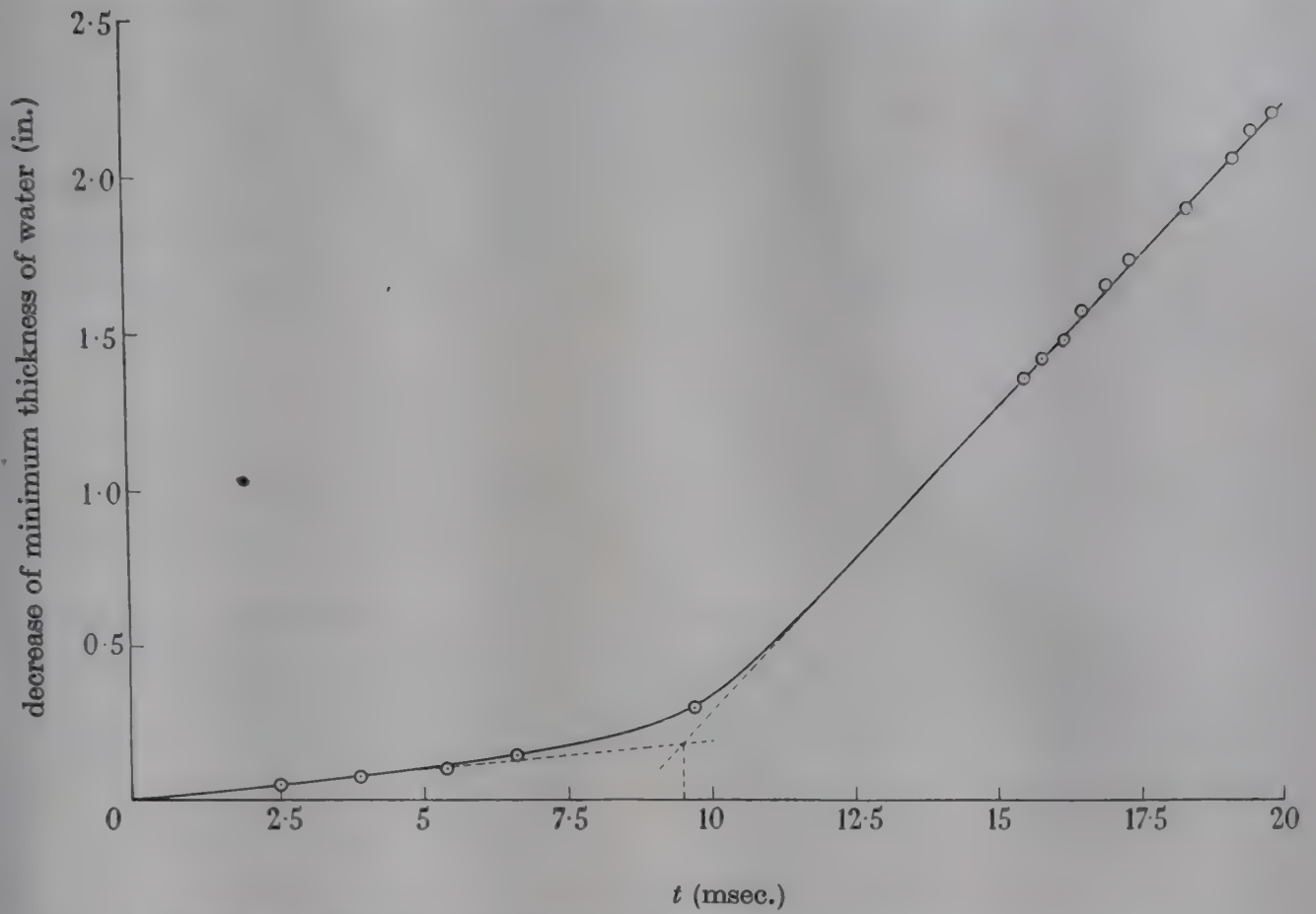


FIGURE 19. Penetration of water by air column. Acceleration of water = $118.5g$; initial depth of water = 5.01 in.; radius of tip of column = 0.80 in.

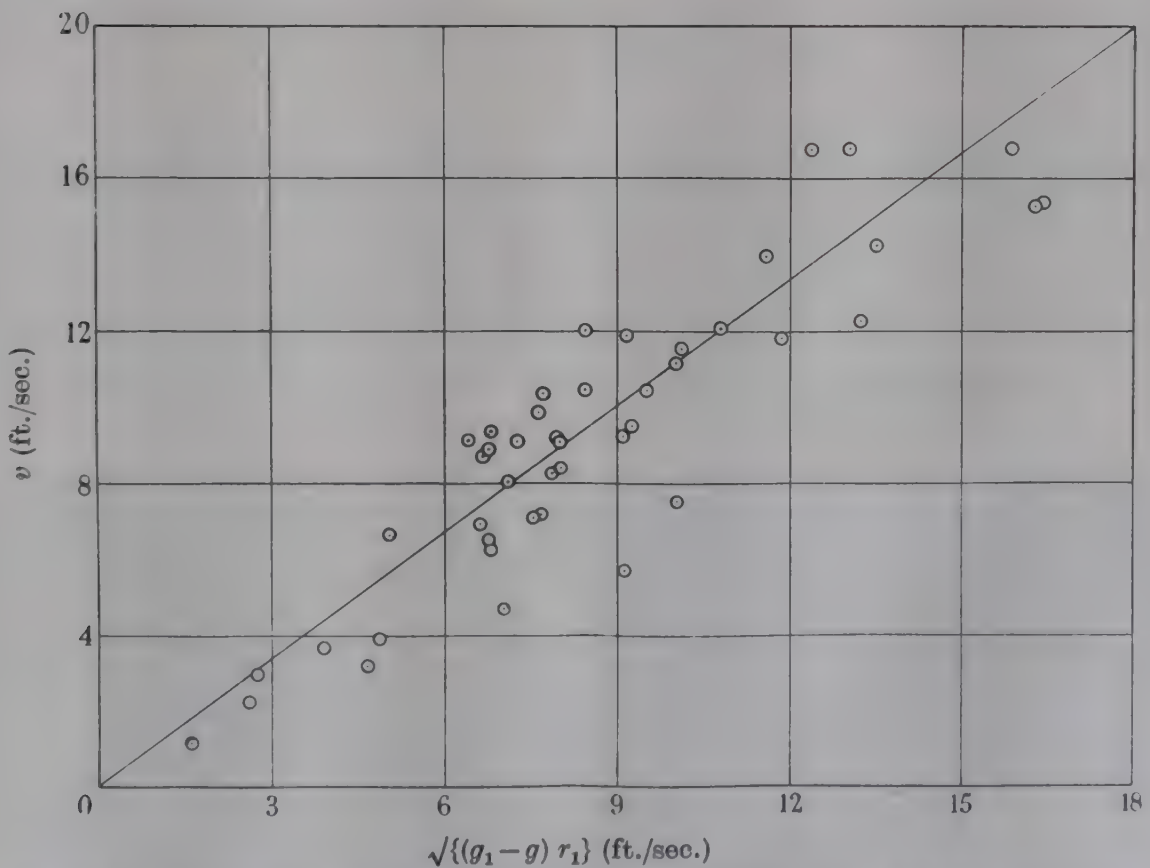


FIGURE 20. Penetration velocities compared with equation (4), $v = 1.11 \sqrt{\{(g_1 - g) r_1\}}$.

In figure 19 it is seen that the time of start of the uniform velocity of penetration can be taken as $t = 9\frac{1}{2}$ msec. From this, the shape factor of the unstable disturbance at that time (η_{c_2}/λ) can be determined. This was calculated wherever possible, and the values obtained range between 0.05 and 2.26 with a mean value of 0.75. There is no systematic variation in the value of η_{c_2}/λ with the acceleration or any of the initial conditions of the experiment. The large range of the values probably arises in the cross-reference to all the derived data and curves for each experiment necessary to derive η_{c_2}/λ .

5. CONCLUSIONS

It is concluded that the instability is made up of the following phases:

- (1) an exponential increase in amplitude as given by the first-order theory until the amplitude is about 0.4λ ;
- (2) a transition stage during which the amplitude increases from 0.4 to $\frac{3}{4}\lambda$ and the surface disturbance changes to the form of round-ended columns of air penetrating into the liquid which forms narrow upstanding columns in the interstices;
- (3) a final stage of penetration through the liquid of the air columns at a uniform velocity proportional to $\sqrt{(g_1 - g)}$.

REFERENCES

- Davies, R. M. & Taylor, Sir. G. 1950 *Proc. Roy. Soc. A*, **200**, 375.
 Taylor, Sir G. 1950 *Proc. Roy. Soc. A*, **201**, 192 (Part I).

A note on three-dimensional turbulence and evaporation in the lower atmosphere

By D. R. DAVIES, *Sheffield University*

(Communicated by L. Rosenhead, F.R.S.—Received 10 January 1950)

In view of the practical significance in dynamical meteorology of the problem of evaporation from areas of finite lateral extent, it is a matter of fundamental importance to test as fully as possible the applicability of the hypothetical three-dimensional model of turbulence which was introduced by the author in 1947. The present paper describes in detail the manner in which the two-dimensional system developed by O. G. Sutton, K. L. Calder and E. L. Deacon for flow over aerodynamically smooth and rough surfaces may be extended to three-dimensional diffusion of vapour over an *evaporating* area. The agreement obtained between theory and experiment is good at points *over* the area. This agreement indicates that the assumed law, introduced by the author to give the variation of the coefficient of lateral diffusivity with height above the surface, may be used satisfactorily in *evaporation* problems as long as attention is confined to points not too far outside the boundaries of the area. The complicated mathematical relationship previously obtained for the vapour distribution vertically above the down-wind edge of a parabolic strip is reduced to a much simpler one. This serves to bring out explicitly the relationship between the effects of the two- and three-dimensional theoretical systems of turbulent transfer at points on the central axis of the area.

1. INTRODUCTION

A theory has been developed recently by Calder (1949) to account, in adiabatic conditions, for the vertical distribution of vapour down wind of rectangular short grass plane surfaces, contaminated uniformly with liquid and orientated with lateral edges perpendicular to the direction of the mean wind velocity during the sampling periods. Account is taken in this theory of the fact that the flow of the atmosphere over grass surfaces is of the 'aerodynamically rough' type, and the formulae involved for concentration of vapour depend only on the wind profile above the surface up to the approximate depth of the vapour layer. The treatment introduced by Calder gives excellent agreement with experiment if the areas are of sufficient lateral extent to be regarded as of infinite width with regard to the sampling points on the centre lines of the areas. The work has been extended to cover non-adiabatic conditions by Deacon (1949). This theory, however, is two-dimensional, and there is a clear need for an extension to three dimensions. In a recent paper, Sutton (1949*a*) has shown how his statistical theory may be applied to cover the problem of three-dimensional diffusion in aerodynamically rough flow in adiabatic conditions from a continuous point source, but this theory in its present form does not appear to be applicable to the problem of evaporation from a finite area. Consequently some assumption must be made with regard to the lateral diffusive properties of the atmosphere. In a paper by the author (Davies 1947) a hypothetical three-dimensional diffusive system is introduced for use in evaporation problems. In this work the well-known two-dimensional conjugate laws giving the variation of wind speed and vertical diffusivity with height are adopted and a law for lateral diffusivity postulated. Using the notation of this paper, the law is written as

$$K_y = a_y U_1^{1-n} z^m. \quad (1.1)$$

It has been shown by the author (Davies 1950, II) that the use of this expression in problems of atmospheric diffusion of gas or air-borne matter from a continuous point source yields an insufficient degree of diffusion from the source, and in view of this deficiency the correct index in the power law was determined to give the exact comparison with the well-known variation of peak concentration of gas with distance directly down wind of the source. The law obtained was

$$K_y = a_y U_1^{1-n} z^{1-m}. \quad (1.2)$$

Substitution of this expression into the general diffusion equation (Davies 1947, equation (2.2)) leads to mathematical difficulties which have not yet been resolved. However, if equation (1.1) is adopted, the evaporation problem is completely soluble in the case of turbulent transfer of vapour from areas of rectangular and parabolic shape, contaminated uniformly with liquid. Arguments have been put forward by the author (Davies 1950, II) to show that the use of equation (1.1) may be expected to result in a solution near the truth for points *over* an evaporating surface, although for points down wind of the area equation (1.2) must be employed.

It is thus a matter of importance to test as far as possible the validity of the use of equation (1.1) in a three-dimensional theory of evaporation. Experimental and

theoretical values have been previously quoted (Davies 1950, II) from work carried out in the Ministry of Supply, giving the vapour distribution at the down-wind edge of a parabolic strip of short grassland, contaminated as uniformly as possible with aniline. These theoretical results were based on a numerical integration of equation (5.6) of a paper (Davies 1947) by the author, together with a method of computing a so-called 'atmospheric' k as detailed on p. 242 of this (1947) paper; the von Kármán k for wind-tunnel flow was multiplied by a factor dependent on meteorological observations to allow for the scale difference of the eddy effect in the wind tunnel and in the open atmosphere. It has been stressed, however, by Sheppard (1947) that the k concept must be used with care in the lower atmosphere because measurements of the mean value of the vertical eddy velocity, i.e. $\overline{w'}$ are rather uncertain. The treatment put forward by Calder holds an advantage in that the formulae do not depend on measurements of $\overline{w'}$, and in near adiabatic conditions a value of $k = 0.4$ is adopted. The theoretical computation of the vapour distribution over a parabolic area has now been repeated by using the ' z^m law', specified in equation (1.1), together with a new value for the constant a_y , evaluated by analogy with the form taken up by the corresponding a_z coefficient in Calder's development.

It has been noticed that the numerical integrations followed in obtaining the results of previous papers (Davies 1947 and 1950, II) are not necessary, and that the theoretical vapour distribution may be expressed mathematically as a series of functions, closely related to the Incomplete Γ -Functions, tabulated by Pearson (1922). The labour involved in the new calculation is thus very much reduced. It is interesting to note that the first term in this series, when applied to the central axis of the area, embodies the two-dimensional result; the following terms may be considered as corrections to give the lateral effect.

2. SPECIFICATION OF A THREE-DIMENSIONAL MATHEMATICAL MODEL OF TURBULENCE, APPLICABLE TO EVAPORATION OF LIQUID FROM A SHORT GRASS, PLANE, SURFACE

The following approximate wind-profile formula was used by Calder (1949) and Deacon (1949) for fully rough flow, and all the entities concerned are fully discussed in the paper by Calder:

$$\frac{u}{v_*} = q' \left(\frac{z-d}{z_0} \right)^m, \quad (2.1)$$

where d is the 'zero-point displacement', and may be taken to be zero for flow over short grass surfaces, and z_0 , the 'roughness' parameter, to be 0.5 cm., independent of wind velocity. The parameter v_* , the so-called 'shearing-stress velocity', is directly proportional to the wind speed. For a wind speed of 500 cm.sec.⁻¹ at 200 cm., Calder determines for v_* a value of 33.3 cm.sec.⁻¹. The mean wind velocity observed in the particular experiment to be considered was 570 cm.sec.⁻¹ at 200 cm. over a short grass surface, and consequently the appropriate v_* value is 38.0 cm.sec.⁻¹; actually this value is also obtained by substituting $u = 570$ and $z = 200$ in the logarithmic profile

$$\frac{u}{v_*} = \frac{1}{k} \log_e \frac{z}{z_0}.$$

Measurements of wind velocity were taken at 2 m. height and also at 1 m., and as the thickness of the vapour layer considered did not exceed 2 m., the value of m was determined from the ratio

$$\frac{u_2}{u_1} = \left(\frac{z_2}{z_1}\right)^m.$$

Substitution of v_* and m into equation (2.1) then gives the value of q' .

The formula used by Calder to include the roughness effect in the vertical diffusivity is

$$\left. \begin{aligned} K_z &= CU_1 z_1^{-m} z^{1-m}, \\ C &= (z_0^{2m}/q'^2 m). \end{aligned} \right\} \quad (2.2)$$

where

In the theory of the paper by the author (Davies 1947) the formula used for the vertical diffusivity was

$$K_z = a_z U_1^{1-n} z^{1-m},$$

and by comparing the two expressions the appropriate value of a_z to incorporate the roughness effect may be evaluated. Thus the following expression is obtained:

$$a_z = (z_0^{2m}/q'^2 m) U_1^n z_1^{-m}. \quad (2.3)$$

In the absence of any theory to specify a_y it is suggested that computation of K_y may be carried out by analogy with the expression accepted for K_z , i.e. write

$$K_y = CU_1 z_1^{1-3m} z^m, \quad (2.4)$$

where the index in the z_1 term is adjusted to maintain the expression dimensionally correct and in keeping with (2.2). Comparison with equation (1.1) yields

$$a_y = (z_0^{2m}/q'^2 m) U_1^n z_1^{1-3m}. \quad (2.5)$$

This procedure is, of course, equivalent to rendering K_y numerically equal to K_z at the reference height z_1 , usually taken to be 200 cm. As exact experimental details are not known for the ratio K_y/K_z , this appears to be the simplest assumption. It is probable that these diffusivity constants are roughly proportional to the corresponding components of the mean-square values of the eddy velocities, and these are known experimentally (Sutton 1949*b*, p. 16) to be of the same order of magnitude in the first few metres above ground, although the lateral component is known to be rather greater than the vertical component. Hence it may be inferred that, at 2 m. height, the ratio of K_y to K_z given by the above equations may not be far from the truth. At heights less than this, probably equation (2.4) underestimates the magnitude of the lateral diffusivity, but the ultimate test of the applicability of equations (2.4) and (2.5) to *evaporation* problems lies in comparison of the resulting theory with experiment.

TABLE 1. NUMERICAL VALUES OF THE PHYSICAL ENTITIES (CM.SEC. UNITS)
EMPLOYED IN THE CALCULATION

$z_1 = 200$	$U_1 = 570$	$v_* = 38.0$
$m = 0.147$	$n = 0.257$	$p = 0.114$
$k = 0.40$	$d = 0$	$z_0 = 0.50$
$q' = 6.20$	$a_y = 14.26$	$a_z = 0.3385$

3. COMPUTATION OF THE VERTICAL DISTRIBUTION OF VAPOUR AT THE DOWN-WIND EDGE OF A PARABOLIC AREA

The formula used is that given by the author (Davies 1947, equation (5.6)):

$$\frac{\chi}{\chi_0} = 1 - \frac{\int_0^\beta \exp(-\mu \sinh^2 \phi) (\sinh \phi)^{2p-1} d\phi}{\int_0^\infty \exp(-\mu \sinh^2 \phi) (\sinh \phi)^{2p-1} d\phi}, \quad (3.1)$$

where χ_0 is the uniform value of the concentration of vapour on the contaminated area,

$$p = n/(2+n),$$

$$\mu = U_1^n A / a_y z_1^m,$$

and β depends on the position of the point under consideration above the area. A is the parabolic parameter given by the equation, $y^2 = 4Ax$, of the boundary of the evaporating area.

To simplify equation (3.1) write

$$\sinh^2 \phi = v, \quad (3.2)$$

and equation (3.1) becomes

$$\frac{\chi}{\chi_0} = 1 - \frac{\int_0^{v_1 = \sinh^2 \beta} \exp(-\mu v) v^{p-1} (1+v)^{-1} dv}{\int_0^\infty \exp(-\mu v) v^{p-1} (1+v)^{-1} dv}. \quad (3.3)$$

In atmospheric conditions with the particular dimensions of the parabolic area considered v is less than unity at points over the area and consequently (3.3) may be written

$$\frac{\chi}{\chi_0} = 1 - \frac{\int_0^{v_1} \exp(-\mu v) v^{p-1} (1 - \frac{1}{2}v + \frac{3}{8}v^2 - \frac{5}{16}v^3 + \dots) dv}{\int_0^\infty \exp(-\mu v) v^{p-1} (1 - \frac{1}{2}v + \frac{3}{8}v^2 - \frac{5}{16}v^3 + \dots) dv} \quad (3.4)$$

$$= 1 - \frac{\left[\Gamma(\mu v_1, p) - \frac{1}{2\mu} \Gamma(\mu v_1, p+1) + \frac{3}{8\mu^2} \Gamma(\mu v_1, p+2) + \dots \right]}{\left[\Gamma(p) - \frac{1}{2\mu} \Gamma(p+1) + \frac{3}{8\mu^2} \Gamma(p+2) + \dots \right]}, \quad (3.5)$$

where the numerator consists of a series of Incomplete Γ -Functions and the denominator a series of Complete Γ -Functions.

By using the relations of the type $\Gamma(p+1) = p\Gamma(p)$, equation (3.5) reduces to the form

$$\begin{aligned} \frac{\chi}{\chi_0} = 1 - & \left[1 + \frac{p}{2\mu} - \frac{p(p+3)}{8\mu^2} + \frac{p(p+1)(10-p)}{16\mu^3} - \dots \right] \\ & \times \left[I(x, p-1) - \frac{p}{2\mu} I(x, p) + \frac{3p(p+1)}{8\mu^2} I(x, p+1) - \frac{5p(p+1)(p+2)}{16\mu^3} I(x, p+2) \dots \right], \end{aligned} \quad (3.6)$$

where $x = \mu \sinh^2 \beta$ and the $I(x, p)$ functions are as defined by Pearson (1922). Both series in practice converge very rapidly, and the process of numerical integration of equation (3.1) used in deriving previous numerical results (Davies 1950, II, table 1a) are no longer necessary.

At any point in a vertical plane at $x = x_1$, it has been shown by the author (Davies 1947, equation (5.13)) that the iso-concentration lines are given by

$$\frac{y^2}{\cosh^2 \beta} + \frac{z^{1+2m}}{(m + \frac{1}{2})^2 \sinh^2 \beta (a_z/a_y)} = 4Ax_1. \quad (3.7)$$

For a given value of (x_1, y, z) this is a quadratic in $\sinh^2 \beta = X$,

$$(4Ax_1)X^2 + [4Ax_1 - y^2 - z^{1+2m}a_y/(m + \frac{1}{2})^2 a_z]X - z^{1+2m}a_y/(m + \frac{1}{2})^2 a_z = 0. \quad (3.8)$$

The value of $\sinh^2 \beta$ obtained from this equation is then substituted in the series (3.6) and the computation is quickly completed by using Pearson's Incomplete Γ -Function Tables, noting that he tabulates $I(x/\sqrt{(1+p)}, p)$.

4. FORMAL CONNEXION BETWEEN THE TWO-DIMENSIONAL AND THREE-DIMENSIONAL THEORIES

At a point on the central line of the area, i.e. $y = 0$, equation (3.7) becomes

$$\sinh^2 \beta = \frac{z^{1+2m} a_y}{(m + \frac{1}{2})^2 4Ax_1 a_z}. \quad (4.1)$$

The argument x in the $I(x, p)$ functions of equation (3.6) becomes

$$\begin{aligned} x &= \mu \sinh^2 \beta \\ &= \left(\frac{U_1^n A}{a_y z_1^m} \right) \frac{z^{1+2m} a_y}{4Ax_1 (m + \frac{1}{2})^2 a_z} \\ &= \frac{z^{1+2m}}{(a_z z_1^m / U_1^n) (2m + 1)^2 x_1} \\ &= \frac{z^{1+2m}}{C(2m + 1)^2 x_1}, \quad \text{from equation (2.2)}. \end{aligned} \quad (4.2)$$

The formula in two dimensions given by Calder (1949, equation (46)) is

$$\frac{\chi}{\chi_0} = 1 - I\left(\frac{z^{1+2m}}{C(2m + 1)^2 x_1}, p - 1\right). \quad (4.3)$$

If the formulae derived from the three-dimensional and two-dimensional theories are denoted by $(\chi/\chi_0)_3$ and $(\chi/\chi_0)_2$ respectively, the relationship between them along the central line, $y = 0$, is

$$\begin{aligned} \left(\frac{\chi}{\chi_0}\right)_3 &= \left(\frac{\chi}{\chi_0}\right)_2 - \left\{ \left[1 + \frac{p}{2\mu} - \frac{p(p+3)}{8\mu^2} + \dots \right] \right. \\ &\quad \times \left[I(x, p-1) - \frac{p}{2\mu} I(x, p) + \frac{3p(p+1)}{8\mu^2} I(x, p+1) - \dots \right] - I(x, p-1) \Big\}. \end{aligned} \quad (4.4)$$

In the particular circumstances of the experiment considered the effect of the corrective terms is very small along the central axis. At points away from the central axis the difference between $(\chi/\chi_0)_3$ and $(\chi/\chi_0)_2$ is *very considerable* and *increases* as the sampling point is moved out laterally.

5. COMPARISON OF THEORY AND EXPERIMENT

The experimental results quoted in table 1 of a paper by the author (Davies 1950, II) are detailed in table 2 below in order to compare with the theoretical results computed by the extension of Calder's treatment to three dimensions. It was noted in the computation that it was rarely necessary to make use of as many as four terms in the series of I functions involved. Direct comparison is not possible since theory gives the concentration $\chi(x, y, z)$ at any point above the area in terms of χ_0 , the ground concentration on the evaporating area. As χ_0 is not known, comparison is made by expressing both theoretical and observed concentrations in terms of the concentration at $z = 1$ in. on the central line $y = 0$, i.e. the maximum observed concentration.

TABLE 2. THEORETICAL AND EXPERIMENTAL RESULTS GIVING MEAN VALUES OF RELATIVE CONCENTRATIONS AT VARIOUS POINTS IN A VERTICAL PLANE ABOVE THE DOWN-WIND EDGE OF A PARABOLIC EVAPORATING AREA. THE AREA WAS 10 YARDS IN LENGTH AND 10 YARDS IN WIDTH ACROSS THE DOWN-WIND EDGE, WITH AXIS DIRECTED ALONG MEAN AIR STREAM AND VERTEX POINTING INTO THE WIND

distance from axis (yd.)	height of sampling point above ground level (in.)							
	1	8	24	48	1	8	24	48
	experimental				theoretical			
0	1.00	0.35	0.06	0.01	1.00	0.36	0.09	0.01
2.5	0.68	0.26	0.04	0.004	0.93	0.29	0.06	0.007
5.0	0.15	0.05	0.01	0.003	0.24	0.03	0.01	0.002

The comparison is seen to be very good and to be an improvement on the values, obtained in table 1 (a) of the author's previous paper (Davies 1950, II), which were based on the concept of an 'atmospheric k '. Even on the outside edge of the contaminated surface, i.e. at a distance of 5 yd. from the axis, the comparison shows that the suggested diffusive laws produce a theoretical distribution of vapour very near to the observed. At a sampling point situated at a distance of 7.5 yd. from the axis of the area a small quantity of vapour was measured but the theoretical values were vanishingly small in comparison. This is probably due to the fact, as previously indicated (Davies 1950, II), that the ' z^m law' yields a point source giving insufficient diffusion except for points near to the source, and hence for points well outside the area the present theory would not be expected to provide a good comparison with experiment.

6. DISCUSSION

The two-dimensional theory, involving the concept of aerodynamic smooth and rough flow, developed recently by K. L. Calder, has been extended to three dimensions, for *evaporation* problems, by introducing a hypothetical law of lateral diffus-

ivity. The resulting diffusive system has been worked out in detail for flow over a short grass surface. It is observed that the application of this system leads to good agreement of theory and experiment with regard to the vapour distribution at all points over a finite short grass surface, bounded by a parabolic curve and contaminated uniformly with liquid. It appears probable that this system of equations may be applied with equal success to surfaces of other shapes, such as the finite rectangle problem discussed by the author (Davies 1950, I), as long as attention is confined to the region over the surface. The extension of the work to the region down-wind of the surface is being investigated.

I am indebted to the Chief Scientist, Ministry of Supply, for permission to publish the experimental figures of tables 1 and 2 which appeared in a Ministry of Supply Report.

REFERENCES

- Calder, K. L. 1949 *Quart. J. Mech. Appl. Math.* **2**, 153.
 Davies, D. R. 1947 *Proc. Roy. Soc. A*, **190**, 232.
 Davies, D. R. 1950 (I and II) *Quart. J. Mech. Appl. Math.* (in the Press).
 Deacon, E. L. 1949 *Quart. J. R. Met. Soc.* **75**, 89.
 Pearson, K. 1922 *Tables of the incomplete gamma function*. London: H.M.S.O.
 Sheppard, P. A. 1947 *Proc. Roy. Soc. A*, **188**, 208.
 Sutton, O. G. 1949*a* *Quart. J. R. Met. Soc.* **75**, 335.
 Sutton, O. G. 1949*b* *Atmospheric turbulence*. London: Methuen.

Resistivity of pure metals at low temperatures

I. The alkali metals

BY D. K. C. MACDONALD AND K. MENDELSSOHN

Clarendon Laboratory, University of Oxford

(Communicated by F. E. Simon, F.R.S.—Received 16 January 1950)

A technique is described for measuring resistance in the temperature interval between 1.5 and 20° K continuously. The resistivity of the alkali metals has been determined in this temperature region, and the results have been discussed in comparison with theory. The resistivity of sodium has been found to agree closely with theory, while in the case of the other metals the hypothesis of quasi-free conduction electrons does not seem to be fully justified. Certain observations suggest that small impurities may give rise to anomalies other than the phenomenon of 'residual' resistance.

1. INTRODUCTION

The modern analysis of the behaviour of the electrical resistance rests essentially on three separate fundamental features:

(1) There is first the nature of the ionic lattice-electron interaction; that is to say, the characteristics of the electric field due to the crystal and the degree of binding of the valence electrons. The information required is summarized more or less

adequately by a specification of the Brillouin zone structure together with a knowledge of the boundaries of the Fermi energy surface, in so far as it may be known. This essentially classifies the intrinsic conductivity of the metal irrespective of thermal vibration or lattice faults.

(2) Secondly, the electron scattering due to the thermal vibration of the lattice, which clearly must be closely linked with the entropy of the latter. The quantitative relation of the scattering to be expected is given by the well-known Grüneisen-Bloch formula for the resistivity of an ideal metal, namely,

$$\rho \sim \left(\frac{T}{\Theta}\right)^5 \int_0^{\Theta/T} \frac{x^5 dx}{(e^x - 1)(1 - e^{-x})}, \quad (1)$$

where T = absolute temperature and Θ = Debye temperature of the lattice. Certain fundamental assumptions underlie this formula as, for instance, that of a quasi-free electron gas of isotropic characteristics. In so far as these are met, equation (1) should then give universally the temperature dependent component of the electrical resistance.

(3) Finally, scatter due to physical or chemical irregularities of the lattice have to be taken into account. Their effect is expressed through the empirical rule of Matthiessen as an additional resistance which is independent of temperature. This 'residual' resistance can be interpreted theoretically if one assumes essentially that the collisions with the thermal Debye waves and with the fault centres are independent of one another.

For very low temperatures ($T \ll \Theta$) the integral in (1) yields a constant value, and hence

$$\rho \sim T^5, \quad (2)$$

as first derived by Bloch (1929). For high temperatures, on the other hand, one sees readily that, since the integral approximates to $\frac{1}{4}(\Theta/T)^4$,

$$\rho \sim T. \quad (3)$$

In general, then, progressive deviation from (3) with falling temperature indicates the onset of quantization phenomena. It is not generally appreciated, however, that in this field significant deviations from the classical conditions, yielding (3), can only be expected to arise for temperatures *well below* Θ . This becomes clear when one compares the variation of classical and quantal entropy of a regular crystal (see figure 1). In order to study significantly the temperature variation of resistivity in comparison with (1) it is thus necessary to work at temperatures below $\sim 0.4\Theta$. For most metals, this requires observation below the temperature of liquid air, and for many below the temperature of liquid hydrogen. On the other hand, the dominance of residual resistance at very low temperatures greatly decreases the value of measurements at liquid helium temperatures in this respect.

The great number of existing observations (cf. Meissner & Voigt 1930) carried out at isolated fixed temperatures (liquid air, liquid hydrogen, liquid helium) is therefore quite inadequate for the comparison with theory envisaged above. Only one series of continuous measurements (de Haas & van den Berg 1936) on a number of metals exists so far. It was therefore decided to begin a systematic determination

of the resistance of all metals in dependence on temperature, particularly in the region between 4 and 20° K. A short note on some of the results has been published recently (MacDonald & Mendelssohn 1948).

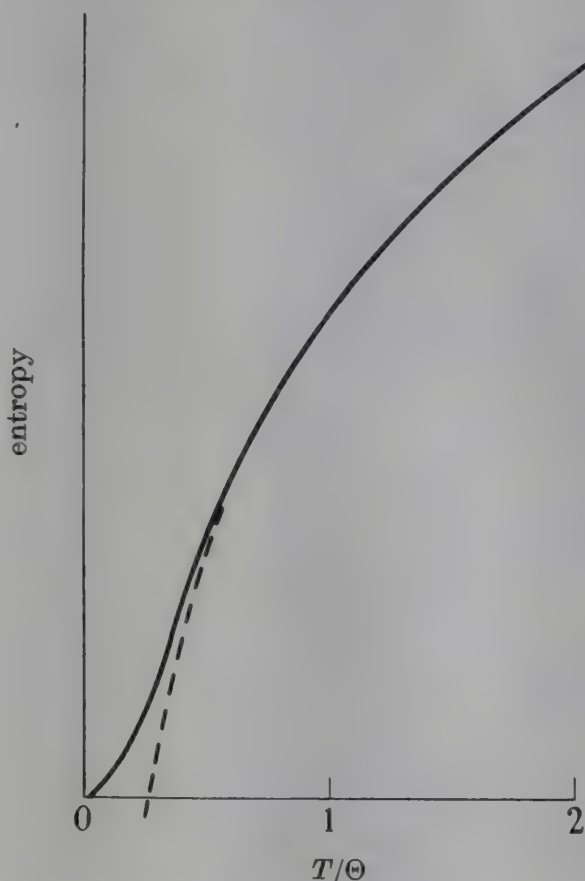


FIGURE 1. Entropy of crystal lattice: —, Debye theory; — — —, classical theory.

2. EXPERIMENTAL PROCEDURE

It is well known that when investigating physical properties in relation to temperature an awkward gap occurs between 4 and 14° K. Above the latter temperature liquid hydrogen boiling under reduced pressure provides a convenient temperature bath; below the former, liquid helium is available. Although solid hydrogen may be pumped down to 10° K or even somewhat lower, considerable difficulties arise owing to its poor thermal conductivity. The Simon desorption method, using helium adsorbed on charcoal, may be used to cover this region, and the Leiden workers have used this method in determining some electrical conductivities over the range 1 to 20° K.

We have, however, been able to use successfully a Simon expansion helium liquefier to cover this whole range continuously with most satisfactory consistency and accuracy, and, in fact, also in a number of cases to cover the range from as high as 90° K (liquid oxygen) with fair accuracy.

In this particular liquefier, devised by Daunt & Mendelssohn (1938), access is provided to the liquid helium container by a long thick-walled tube extending to the top of the whole apparatus (see figure 2), which is normally sealed with a high-pressure nut during the liquefaction process. For this investigation this tube was

used to enable a gas thermometer bulb, adjacent to which the specimen was placed, to be inserted for the period of the experiment. This bulb communicates through a fine capillary with an accurate large-scale vacuum gauge mounted on the top of

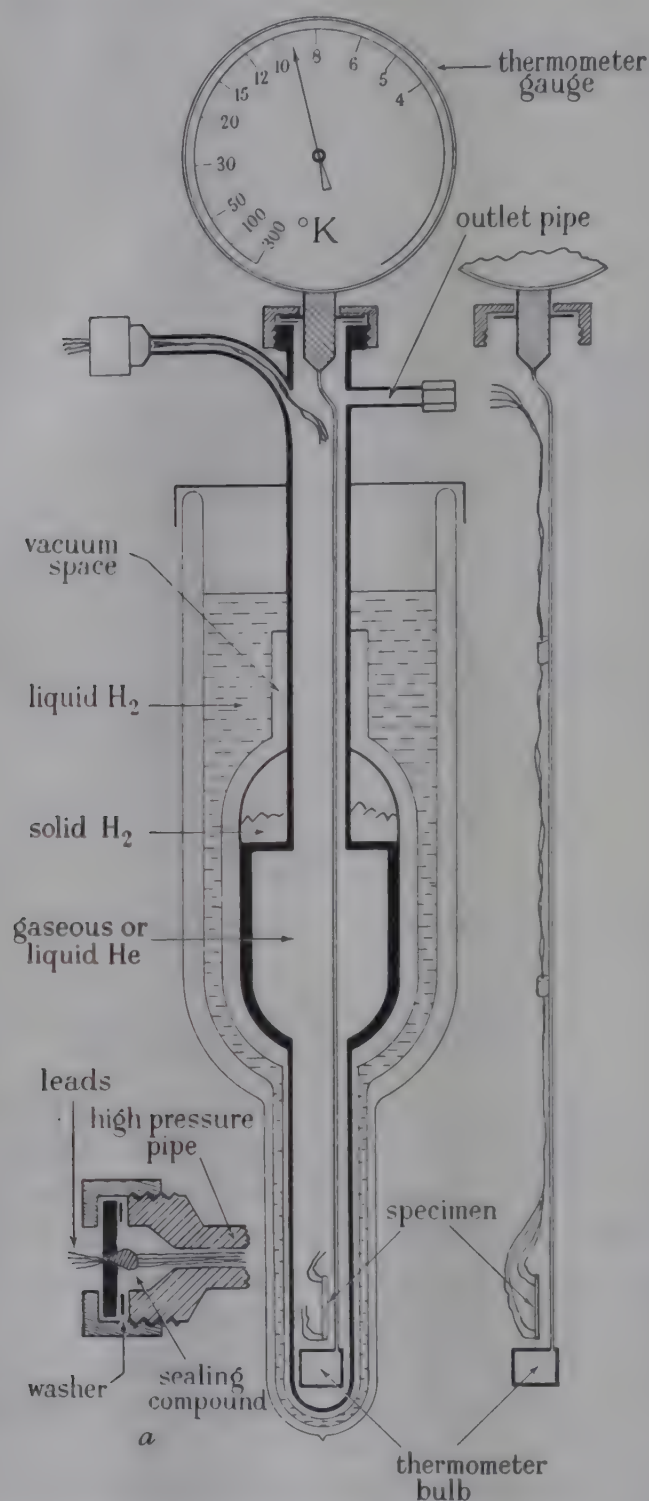


FIGURE 2. Cryostat (not to scale).

the apparatus; the whole is filled with very pure helium gas at about 1 atm. at room temperature. Suitable choice of the ratio of bulb volume to that of the gauge (cf. Mendelssohn 1931) enables one to arrange for the scale to be very 'open' over a more or less arbitrary temperature interval. In our case about 30 cm. of scale interval corresponded to the range 4 to 20° K.

Each specimen was provided with separate current and potential leads, and these were brought out through a high-pressure seal. Considerable effort went towards providing an effective simple seal. Fair success was achieved with a seal of the type shown in figure 2*a*, where the insulated leads emerge through a conical aperture terminating in a very small hole in a brass boss and are sealed in with hot sealing-wax. Although Bridgman (1931) advocates such a seal for pressures of 100 or 200 atm. considerable difficulty was met with in trying to make the wax adhere satisfactorily to the metal boss. Hard Bakelite or similar materials for the boss proved unsatisfactory owing to distortion under pressure. Ultimately, the best solution was found in the use of small commercial Kovar glass-to-metal seals soldered to a brass cap. The tubular type of seal was used, since this allowed the potential leads to come straight through without the intervention of a thermo-electric junction with the Kovar metal.

Finally, an outlet equipped with a small high-pressure nut was provided so that after liquefaction of the helium the nut could be removed under atmospheric pressure and connexion made to a vacuum pumping system to reduce the temperature still further.

The essential feature of the method then is that, particularly below 20° K, the highly compressed helium gas in the liquefier, in which the specimen and thermometer capsule are situated, provides a large heat capacity, practically equivalent to a liquid bath. It should be borne in mind that the density of the helium gas before expansion may well exceed that of the final liquid helium. Furthermore, since the gas is enclosed in a vessel of high heat conductivity, excellent thermal equilibrium is to be expected. During the first part of the run, from 20 to ~12° K, constant temperatures may readily be achieved by control of the pumping-off rate of the hydrogen condensed inside the expansion liquefier, and during the latter part, ~12 to 4° K, by adjustment of the expansion rate of the helium gas. Adjustment could be made by manual control of a needle valve in the expansion line to give a temperature constancy of 0.1° K or better. By suddenly changing the trend of the temperature drift and simultaneously observing the resistance of the specimen it could be established that the latter followed the indications of the thermometer without appreciable delay. Finally, after production of liquid helium, connexion may be made to a constant-pressure pumping device (cf. Mendelssohn 1936) and constant temperatures down to ~1.5° K attained.

3. ELECTRICAL MEASUREMENTS

A galvanometer amplifier using simple selenium photo-voltaic cells was designed to measure the small potentials developed. Galvanometer amplifiers have been used before in this field (e.g. Milner 1937; Kapitza & Milner 1937), but in these cases the photo-cell amplification has been used solely to improve the sensitivity of null detection in a bridge circuit. Systematic or random drift is generally rather troublesome in such systems, and although a relatively straightforward optical circuit is intrinsically capable of providing an enormous magnification, only a very small fraction can be usefully employed under these circumstances. The incorpora-

tion of negative feed-back in such a galvanometer amplifier enables the 'unwanted' amplification to perform valuable service in a number of ways. By returning a large fraction of the output of the photo-cell in opposition to the input voltage rather than 'neglecting' it in a step-down potentiometer one can provide a greatly enhanced stability of the overall system, a considerably reduced response time and a high linearity while still leaving ample net gain for the purpose in hand. Preston (1946) used such a system with the feed-back arranged in parallel with the input terminals, which has the effect of greatly lowering the effective input impedance to the system—a valuable feature when it is desired to record from a low-resistance source such as a thermocouple, essentially a current-generating device. In our case, it was desired to record directly from the potential leads of the specimens in which

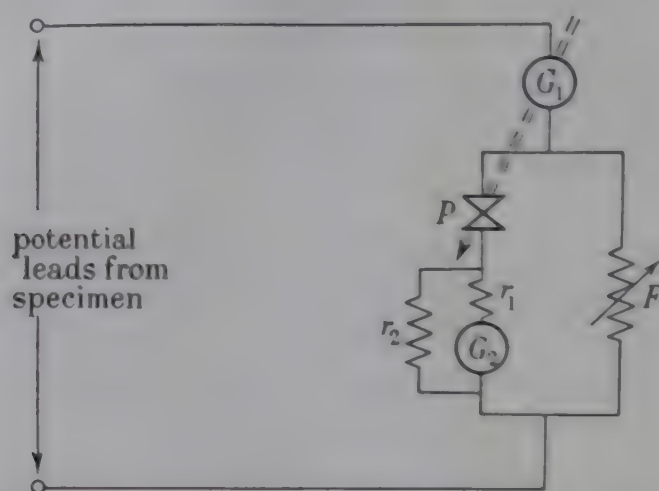


FIGURE 3. Basic circuit of galvanometer amplifier. G_1 , amplifying galvanometer; G_2 , recording galvanometer; r_1, r_2 , local shunt for G_2 ; P , split photo-cell; F , feed-back control resistance. $\angle =$, optical beam reflected from G_1 mirror to P .

it is necessary, to avoid the effect of contact resistance, to draw practically *no* current; that is, a device with a very high input impedance is called for. This is achieved by modifying the system (MacDonald 1947; cf. Frankenhaeuser & MacDonald 1949) so that the feed-back is generated in series with the input; this results in an input resistance at d.c. of at least several thousands of ohms which may certainly be neglected in comparison with any contact resistance likely to be encountered in such work.

The circuit diagram of the galvanometer amplifier is shown in figure 3. The galvanometer, G_1 , was a Tinsley high-sensitivity galvanometer ($R \approx 30 \Omega$, $\sim 1000 \text{ mm./}\mu\text{A}$), while the recording galvanometer, G_2 , was a Cambridge high-sensitivity short-period galvanometer ($R \approx 450 \Omega$, $\sim 1500 \text{ mm./}\mu\text{A}$); a scale length of 1 m. was used with a 3 m. optical throw with an image which could be read to $\sim \frac{1}{2} \text{ mm}$. An upper limit of relative accuracy $\sim 1/1-2000$ was thus available. This appears entirely adequate for this type of investigation. Furthermore, when a pure metal specimen at very low temperatures provides a total potential difference of only $\sim 5-10 \mu\text{V}$ with a measuring current $\sim 100 \text{ mA}$, there is no advantage in seeking for a sensitivity greater than $\sim 10^{-8} \text{ V}$. The greatest advantage of a direct-reading system for these high sensitivities lies in the fact that no variable contacts exist in the potential

measuring circuits. The only switch that occurs is in the *current* network where thermal and contact e.m.f.'s are of no consequence. In a potentiometer or other nulling circuit such e.m.f.'s are *in principle* compensated for by commuting the circuits. Should the disturbing e.m.f. vary too rapidly commuting must fail to achieve the object. With a direct-reading system, however, the effect of any disturbing e.m.f. is continually visible and may be readily corrected for. The method used by Kapitza & Milner (1937) of overcoming this difficulty in a potentiometer by rapid-action switches was tried but proved unsatisfactory.

The problem of stability as related to the choice of a particular photo-cell lies outside the scope of this paper, and is being discussed elsewhere (MacDonald 1950).

4. PREPARATION OF SPECIMENS

In view of the great chemical activity of the alkali metals, particular care had to be taken in the preparation of the specimens. The four metals (Na, K, Rb, Cs) were prepared in soft glass capillary moulds with platinum electrodes sealed in. The glass apparatus was thoroughly cleaned and washed out with distilled water. It was then pumped under hard vacuum for an hour or two whilst being heated at $\sim 150^\circ \text{C}$. Pure helium to 1 atm. was then admitted; the glass capsule containing the metal was broken and immediately inserted through a cone and socket joint. In the case of rubidium and caesium the capsules were severed under benzene previously dried thoroughly with sodium wire.

The capsule was inserted, and the system was rapidly re-evacuated before the film of benzene had evaporated from the cut surface. Formation of oxide was by this means practically entirely obviated. After melting, the metal was forced into the capillary by admitting pure helium gas at 1 atm. The specimen was then allowed to cool very slowly under this pressure. It was rare that any sample prepared under these circumstances did not then form a continuous specimen, whose resistance agreed satisfactorily with that computed from the geometry of a solid wire. When cool, the specimen tubes were cracked off (under benzene where necessary), sealed over with Plasticine and covered finally with a suitable sealing compound (Durofix). When this has hardened a satisfactorily air-tight seal is formed. To avoid the risk that small cracks, which might develop in the seal on cooling to helium temperatures, would vitiate a future experiment, the specimens were kept between runs in dried petrol ether.

In the case of lithium, this technique cannot be employed because as soon as the lithium melts (m.p. 186°C) and comes into contact with ordinary glass a relatively violent reaction occurs, the glass cracking almost immediately. Lithium glass,* on the other hand, proved to be extremely soft and very difficult to work, and it has not yet been possible to make suitable capillary specimens. The following technique was finally adopted for lithium. The metal was melted in a stainless steel crucible, in a closed atmosphere of pure helium. A glass capillary mould with open end was

* We are indebted to the British Thomson Houston Co. Ltd for kindly undertaking to make this glass for us.

then quickly inserted into the molten metal after the surface oxide had been skimmed off. A tap connecting the upper end of the capillary with a vacuum bottle was then immediately opened. Lithium then rushed up into the mould and froze almost immediately without having had time to react appreciably. Suitable specimens resulted on about one occasion in three; difficulty was always found, however, with the upper contacts, as the metal normally only covered the lower electrodes, having frozen in the capillary before reaching the upper contacts in sufficient quantity. Fine needles were therefore inserted by hand into the metal to provide the remaining pair of electrodes, and proved quite satisfactory. The specimen was then sealed off as before.

5. EXPERIMENTAL RESULTS

(1) Sodium

The results on sodium are discussed first, since their interpretation appears simplest from the theoretical standpoint. Metal of high purity was kindly given to us by Messrs British Thomson Houston Co. Ltd. It was supplied after triple distillation in hard vacuum of commercial sodium in sealed boro-silicate glass tubes.

Spectrographic analysis* gave the following result:

	not detected
(parts in 10^6)	(parts in 10^6)
B: 2000	Cs: < 200
	Hg: < 200
Si: 2000	Sn: < 10
	Pb: < 5
K: 500	Rb: < 5
	Li: < 1

It was later confirmed that the boron and silicon impurities had come from the glass in the preparation of the specimen for analysis, and the intrinsic purity of the metal as used by us may therefore be taken as 99.95 %.

Specimens Na 1 and Na 2 of $\sim 150\mu$ diameter exhibited entirely similar behaviour. Some forty observations were made on specimen Na 1 in liquid helium at temperatures from 4.2 to $\sim 1.5^\circ$ K, and the resistance was found strictly constant within 0.4 % with no systematic variation. Subsequent experiments on Na 1 and Na 2 were made covering the range from ~ 20 to 4.2° K. The resulting curves were essentially coincident after the residual resistance had been subtracted, and the data for Na 2 appear in figure 4. A logarithmic plot of the data is given in figure 5. If we assume that the ideal resistance is representable by an expression $R \sim T^n$, then the slope of the logarithmic plot gives the index n .

The value $n \approx 4.85$ is found for the specimen Na 2. This index agrees very closely with the theoretical value 4.89 expected for the temperature range 8– 20° K as computed from equation (1). The value of exactly 5 for the index (cf. equation (2)) will not of course be reached except for vanishingly low temperatures, and it is in fact impossible to deduce reliable values for the ideal resistance of sodium below

* This analysis as well as that of the caesium metal was kindly furnished by A.E.R.E. Harwell.

about 8°K because of the extreme steepness of the curve. The ratio of the resistance at 4.2°K to that at 0°C or at room temperature is frequently quoted as a rough measure of the purity (chemical and physical) of the material and for Na $2R_{4.2^\circ\text{K}}/R_{290^\circ\text{K}} \approx 7 \times 10^{-4}$. It is immediately obvious that this ratio may be used validly for comparison of the purity of various specimens of the *same* metal, since

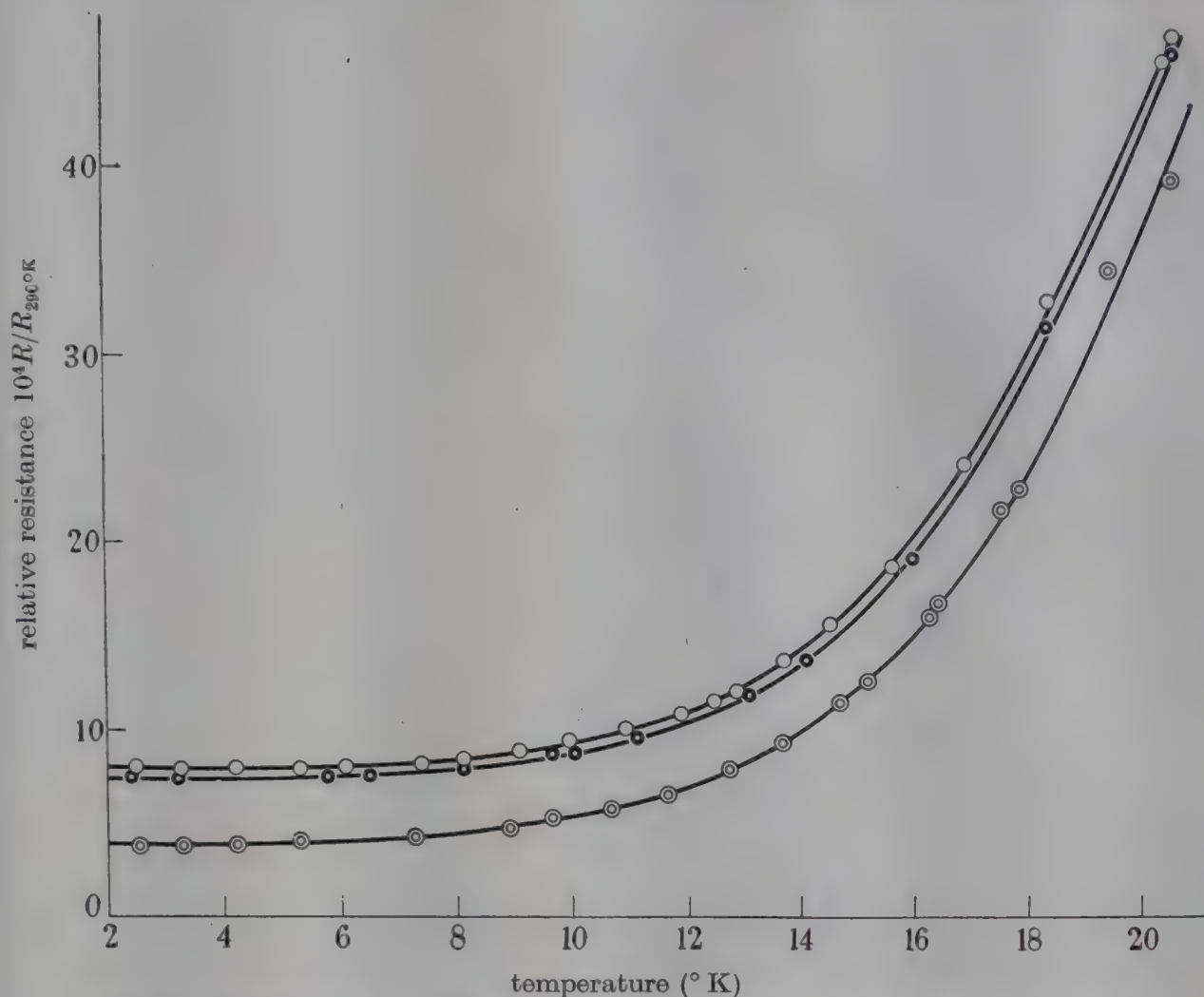


FIGURE 4. Resistance of sodium below 20°K . Specimens: $\circ$, Na2; $\odot$, Na3; $\circ$, Na4.

the resistance at 4°K is generally dominantly 'residual' resistance. More strictly, of course, one should use the value of resistance as $T \rightarrow 0^\circ\text{K}$. On the other hand, the value at 0°C acts effectively as a normalizing factor to remove the necessity of including the geometry of the particular specimen. This again is valid, since for all reasonably pure metals the scattering at such a high temperature is effectively purely thermal in origin and thus independent of purity.

The matter is less obvious when one tries to compare different metals. One may, on the one hand, turn to the theoretical conductivity expressions and try thus to determine the significance of the ratio for different metals. Unfortunately, the expressions then involve such factors as the effective number of free electrons per atom which are by no means certain *a priori* theoretically. One therefore approaches the problem more directly. At 0°K the mean free path is given by

$$l_{0^\circ\text{K}} = \frac{1}{n_i A_i} = \frac{\Omega_0}{x_i A_i},$$

where n_i = number of impurity atoms/cm.³, x_i = number of impurity atoms per cent of parent metal, Ω_0 = atomic volume of parent metal, A_i = scattering cross-section of impurity (and where, if more than one kind of impurity is present, $x_i A_i$ is to be replaced by $\Sigma x_i A_i$). Therefore

$$x_i A_i = [r] \frac{\Omega_0}{l_0^\circ \text{C}},$$

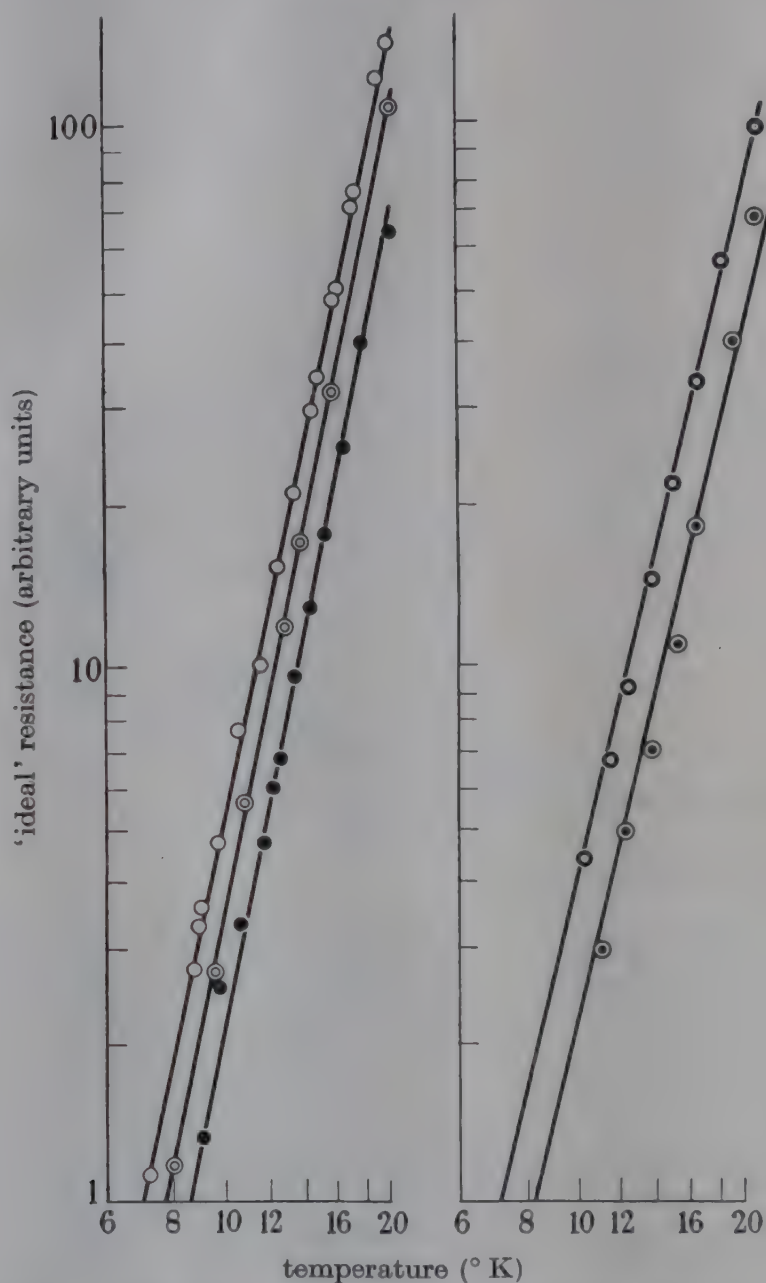


FIGURE 5. Logarithmic plots of the results for sodium and lithium.
 ⊙, Na2; ○, Na3; ●, Na4; ◌, Li1; ⊖, Li3.

where $[r]$ is the ratio of resistance $R_{0^\circ \text{K}}/R_{0^\circ \text{C}}$. Thus it is clear that the ratio $\Omega_0/l_0^\circ \text{C}$ is to be regarded as characteristic of each metal in determining its effective impurity content $x_i A_i$. We will then accept the experimental values of $l_0^\circ \text{C}$ (cf. Mott & Jones 1936, p. 268) and can thus calculate the appropriate factor for each metal. For the alkali metals in particular we find the following values for the parameter $\Omega_0/l_0^\circ \text{C}$, normalized to unity for sodium:

metal	Li	Na	K	Rb	Cs
$\Omega_0/l_0^\circ \text{C}$	1.6	1	1.74	3.6	6.2

TABLE 1. SODIUM: $\Theta = 202^\circ \text{ K}$

$T (^\circ \text{ K})$	$r_{\text{calc.}}$	$r_{\text{obs.}}$		
		(1)	(2)	(3)
273.2	1.0000	1.0000	1.0000	1.0000
170.87	0.5672	—	0.5672	—
108.72	0.3135	—	0.3168	—
90.0	0.2600	—	—	0.2420
87.8	0.2279	0.2279	—	—
77.6	0.1860	0.1849	—	—
56.77	0.1022	—	0.1055	—
20.4	0.00327	0.0034	—	0.00326
15.95	0.00100	—	—	0.00098
14.1	0.00055	—	—	0.00051
13.1	0.00038 ₄	—	—	0.00036 ₅
11.05	0.00015 ₄	—	—	0.00017 ₃
9.65	0.00009 ₃	—	—	0.00010 ₂
8.1	0.00004	—	—	0.00005
4.2	0.00000	—	—	0.00000

Key: (1) Meissner & Voigt (1930); (2) Woltjer & Kamerlingh Onnes (1924); (3) present work.

TABLE 2. LITHIUM: $\Theta = 363^\circ \text{ K}$

$T (^\circ \text{ K})$	$r_{\text{calc.}}$	$r_{\text{obs.}}$		
		(1)	(2)	(3)
374.5	1.433	1.443	—	—
329.7	1.243	1.247	—	—
273	1.000	1.0000	1.0000	1.0000
90.9	0.1692	—	0.1621	—
90	0.166	—	—	0.171
86.3	0.1492	—	0.1464	—
80.1	0.1236	—	0.1252	—
77.7	0.1139	—	0.1169	—
20.4	0.0004	—	0.0013	0.0013 ₃
17.8	0.0002 ₂	—	—	0.0007 ₆
16.05	0.00015	—	—	0.0004 ₅
14.8	0.00011	—	—	0.00029
13.55	0.00008	—	—	0.0002 ₁
12.3	0.00004	—	—	0.0001 ₄
11.4	0.00002 ₄	—	—	0.00009

Key: (1) Meissner (1920); (2) Meissner & Voigt (1930); (3) present work.

TABLE 3. LITHIUM: $\Theta = 268^\circ \text{ K}$

$T (^\circ \text{ K})$	$r_{\text{calc.}}$	$r_{\text{obs.}}$	
		(2)	(3)
273	1.0000	1.0000	1.0000
90.9	0.222	0.1621	—
80.1	0.174	0.1252	—
20.4	0.0013 ₃	0.0013	0.0013 ₃
14.8	0.0002 ₅	—	0.0002 ₉
12.3	0.0001 ₂	—	0.0001 ₄

Key: as for table 2.

Thus it is clear that in fact the ratio $[r]$ alone can only be regarded as a very rough guide to provide comparative purities. In our case, however, the inclusion of $\Omega_0/l_0^\circ\text{C}$ merely serves to emphasize the remarkable purity attainable in the case of sodium.

A specimen (Na3) was kept after manufacture for many months at room temperature. The resistance ratio (see figure 4) had now dropped to $\sim 3.8 \times 10^{-4}$ which presumably indicates that a considerable degree of annealing had occurred, probably leading to the formation of a monocrystalline sample. This then appears to be the purest sodium specimen recorded. The power index deduced was practically identical ($n \approx 4.83$) with that for the earlier specimens, thus providing confirmation of Matthiessen's rule. A further sample (Na4) was made in a very wide capillary (~ 1 mm.). Sufficient measurement sensitivity was achieved by winding about 30 cm. of the capillary in spiral form. This specimen of large diameter was examined in order to eliminate the possibility that a limitation of electron mean free path by the physical size of the specimen might modify the observed data. The results on Na4 (see figures 4 and 5) agree closely with those on the previous specimens, yielding an index 4.86.

A more detailed comparison with theory may be made by comparing the observed values of resistance with those computed from equation (1) over a wide temperature range. Table 1 collates the data of a number of investigations, including the present paper, covering the range from $\sim 0^\circ\text{C}$ to $\sim 10^\circ\text{K}$. The agreement with the Bloch-Grüneisen theory is seen to be excellent, and one should note particularly that a single value of the parameter Θ (namely, $\Theta = 202^\circ\text{K}$) suffices throughout. It is very interesting to compare this value with that deduced from the temperature variation of the specific heat. The value of the Debye temperature necessary to give reasonable agreement with the experimental data at low temperatures (~ 10 to 60°K) is about 150°K . Simon, however (1926, 1930), has pointed out that the values for sodium and lithium derived from specific-heat data are inconsistent with those obtained from the electrical resistivity and the Nernst-Lindemann formula. He has interpreted therefore, as previously in the case of grey tin, the data as composed of a normal Debye spectrum plus an 'anomalous' transformation obeying a Schottky (1922) function. It seems therefore significant that our work on the electrical conductivity has yielded a much higher value of Θ for sodium, namely, 202°K as postulated by Simon, covering the whole temperature range. It thus appears that the postulated internal excitation is of such a nature as not to affect the electrical conductivity. The resulting Schottky contribution has a characteristic temperature of 95°K . On the other hand, the possibility that the true vibrational spectrum of an alkali metal crystal deviates considerably from a simple Debye spectrum as suggested by more recent theoretical work (e.g. Blackman 1937*a, b*; Kellerman 1940, 1941) cannot be excluded. Support for this view arises from the considerable elastic anisotropy of the alkalis (cf. Fuchs 1936).

Observations were also made on Na2 throughout the range 20 – 4°K of resistance in a uniform transverse magnetic field of 4000 gauss, and no change as great as $\sim 0.25\%$ could be detected.

Thus our results without and with a magnetic field give no indication of the small specific heat anomaly at $\sim 8^\circ\text{K}$ found in sodium by Simon & Pickard (1948).

Since the entropy change in this anomaly was very much less than $R \ln 2$, it was suspected to be of electronic origin and might possibly have shown up in conduction experiments. Equally, we have found in none of our specimens any indication whatsoever of a rise in resistivity with falling temperature below 4°K such as reported by Meissner & Voigt (1930).

(2) Lithium

Metal from three sources was examined. Li 1 was made from metal received through the courtesy of the late Dr R. A. Hull, and it is believed that its original source was Kahlbaum. Li 2 was made from metal supplied by Messrs New Metals and Chemicals Ltd., while the metal for Li 3 was obtained through the courtesy of Dr W. Hume-Rothery. All three samples appear to be of approximately the same overall purity, judged from the residual resistances ($\sim 3.3 \times 10^{-3}$; $\sim 4 \times 10^{-3}$; $\sim 5 \times 10^{-3}$ respectively).

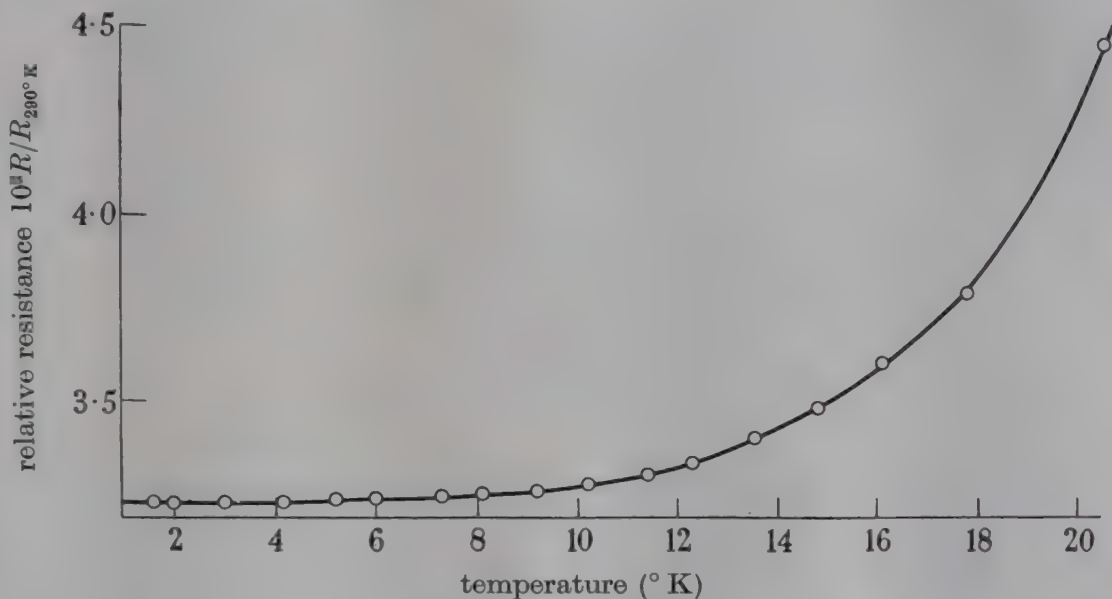


FIGURE 6. Resistance of lithium (specimen Li 1) below 20°K .

Resistance measurements were made on all three specimens, but the results on Li 2 showed appreciable scatter not normally encountered, and therefore only the results from Li 1 and Li 3 were utilized. The resistance was measured for both these specimens down to $\sim 1.8^\circ \text{K}$ and presented no obvious abnormalities; in particular, the resistance was constant below $\sim 4^\circ \text{K}$. The resistance-temperature curve for Li 1 is shown in figure 6. A logarithmic plot for Li 1, after subtraction of residual resistance (figure 5), shows quite definitely an index $n = 4.52$. The data for Li 3 are somewhat less certain, but a line of slope $n = 4.56$ fits the data satisfactorily.

We thus find disagreement with theory in this case, since lithium with a Debye temperature of $380\text{--}400^\circ \text{K}$ would be expected to obey a T^5 law almost rigorously below 20°K . The observed behaviour concurs with the progressive departure from Grüneisen's formula towards low temperatures (20°K) found by previous workers (Meissner 1920; Meissner & Voigt 1930). A value of $\Theta = 363^\circ \text{K}$ (in general satisfactory agreement with specific heat data) is necessary to secure reasonable agreement between room temperature data with those in the region $70\text{--}90^\circ \text{K}$, but

considerable deviation is then already evident at 20°K (see table 2, where the data on Li 1 have been collated with those at higher temperatures of earlier workers). If, on the other hand, we enforce agreement at $\sim 20^\circ\text{K}$ we find a value $\Theta = 268^\circ\text{K}$ is called for. It seems very difficult to assume such a low value, although it must be admitted that the Debye temperature deduced from the specific heat is falling off progressively at low temperatures (at 80°K , $\Theta = 380^\circ\text{K}$, while at 20°K , $\Theta = 340^\circ\text{K}$). It would therefore seem very valuable to have further specific heat data for lower

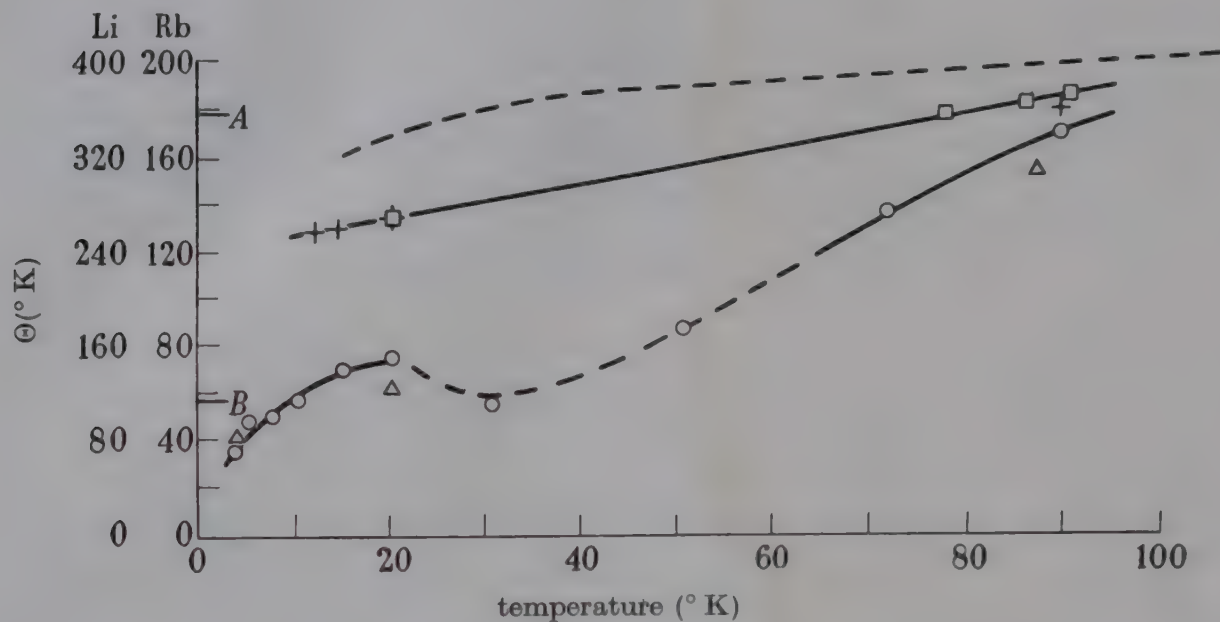


FIGURE 7. Variation of Θ with temperature.

	present work	deduced from data of Meissner & Voigt (1930)
Rb	○	△
Li	+	□

The top dashed line represents Θ for Li as deduced from specific heat data. *A* is the limiting value of Θ for Li deduced from the elastic constants and *B* is the value of Θ for Rb deduced from Lindemann's melting-point formula.

temperatures, and this remark applies also to the other alkali metals yet to be discussed. A selection of resistance values has been calculated in table 3 with $\Theta = 268^\circ\text{K}$ for comparison with the data, and while reasonably good agreement below 20°K is then obtainable, serious discrepancies now emerge in the upper region, as might be expected. The data have also been expressed in the form of a curve representing the variation of Θ with temperature as necessary to procure agreement with Grüneisen's formula and are presented along with data from rubidium in figure 7.

(3) Potassium

A rather detailed investigation was carried out on this metal and seven specimens in all were examined. The sources, treatment and residual resistance ratios are tabulated below, the latter giving a qualitative trend of purity:

specimen	source	treatment	residual resistance ratio
K 1	Hopkins and Williams Commercial; suppliers unable to quote any analysis	melted under vacuum into capillary as received	$\sim 1 \times 10^{-1}$
K 2	G.E.C. ex B.D.H. (quoted 'typical' analysis Na 0.1 % Cl 0.01 % SO ₄ 0.01 % NH ₃ 0.01 % heavy metals 0.002 %)	melted under vacuum into capillary as received	$\sim 5 \times 10^{-2}$
K 3	Dr R. A. Hull (E. Bol- ton King; reported once distilled)	melted under vacuum into capillary as received	$\sim 3 \times 10^{-2}$
K 4	G.E.C. (once distilled by them)	melted under vacuum into capillary as received	$\sim 4.2 \times 10^{-2}$
K 5	as K 3	distilled twice in Pyrex glass by us	$\sim 1.75 \times 10^{-2}$
K 6	as K 3	distilled five times in Pyrex with rejection of final top- fraction distillate	$\sim 1.45 \times 10^{-2}$
K 7	as K 3	remainder of previous sam- ple re-distilled twice in potassium glass (B.T.H.)	$\sim 7.5 \times 10^{-3}$

The first sample examined, when plotted logarithmically over the range 20 to 6° K, gave a law $R \sim T^{3.2}$. This agrees with van den Berg's observation (1938) that a cubic index is valid over the range 14 to 20° K. Difficulty again arises here in essaying comparison with theory. From about 70° K upwards satisfactory agreement is found with Grüneisen's formula if Θ is taken $\sim 163^\circ$ K. On the other hand, a value $\Theta \sim 70^\circ$ K would be called for to procure agreement in the low-temperature region, while the Debye temperature deduced from specific heat data from ~ 14 to 80° K is about 100° K. In view of this conflict of evidence and the general uncertainty in our knowledge of the vibrational spectrum it is unreasonable to be dogmatic in interpreting these results. We can, however, suggest that the upper value of Θ does indicate the basic spectrum as in the case of sodium and lithium on Simon's hypotheses (1926, 1930), and we have then to explain the slow decrease in electrical resistance at the low temperatures.

The resistance curve also showed, however, a slight anomalous hump centred about 12° K and extending roughly from 10 to 14° K. This anomaly had a magnitude only $\sim 1/3$ %, but in view of the high consistency of all the measured points it was felt to be significant. A somewhat purer specimen, K 2, exhibited the anomaly more strongly, and a detailed experiment on the specimen located a cusp of about 1 % magnitude at 13.4° K. In the experiments on K 3 and K 4, again of slightly higher purity, the original anomaly remained, but a second one at $\sim 10^\circ$ K became

evident. The experimental points for K 4 and K 6 appear in figure 8. These anomalies were also confirmed in specimens K 5 and K 6, which showed a further improvement in purity.

It will be noticed that little significant improvement in residual resistance occurred with the repeated distillation of K 6 over K 5. It was presumed that this must be due to collection of impurity atoms from the glass, compensating the improvement by distillation. Messrs British Thomson-Houston Co. Ltd. kindly undertook then to manufacture some potassium glass for us, and this was used for

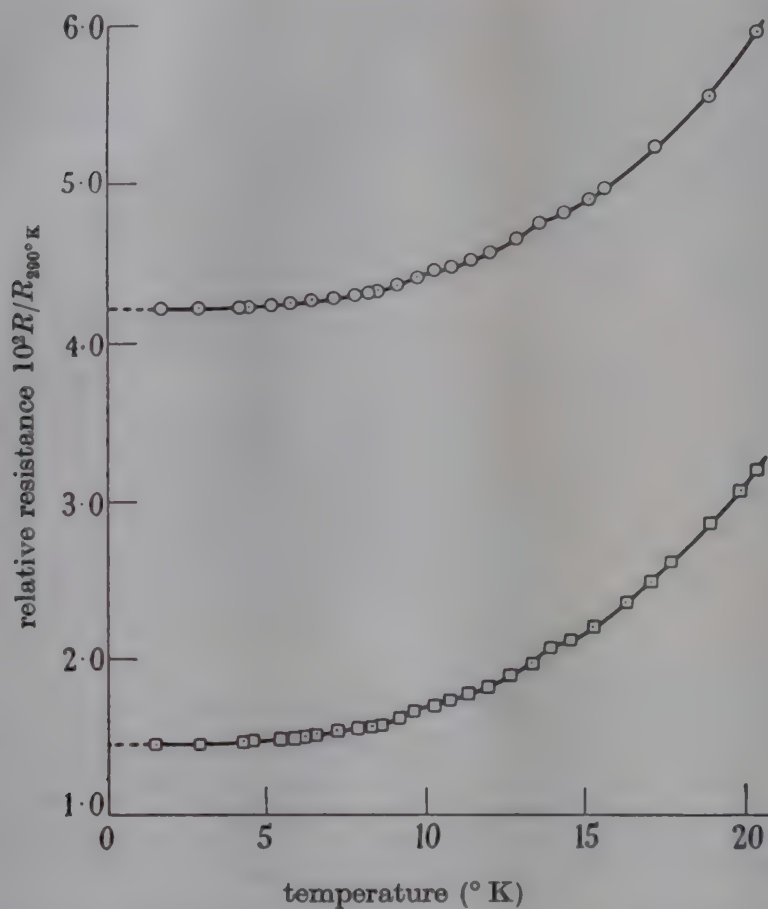


FIGURE 8. Temperature-variation of potassium below 20° K. Specimens: $\odot$, K 4; $\square$, K 6.

distillation to produce K 7. This specimen then showed a satisfactory decrease in residual resistance and also the anomalies had disappeared. We therefore deduce that these resistive anomalies must be characteristic of a relatively small number of impurity atoms in solid solution. On the one hand, the pure metal does not exhibit the phenomenon, while of course any serious amount of impurities produces so much random scattering in any case that any coherent effect due to their presence is entirely masked by a high residual resistance. In view of the later results on barium* it seems very probable that the significant impurity is sodium.† Hard glass was in fact used for the earlier distillations, because of its lower sodium oxide

* In course of publication.

† Cf. also the quoted analysis for K 2.

content than soft glass;* the percentage content of sodium oxide is none the less quite considerable.

It is interesting in relation to these experiments to notice that if van den Berg's data on his first potassium specimen—of the same order of purity as our K 3—are plotted, it is quite possible to fit a curve showing similar cusps, although there is a paucity of data in the relevant region. On the other hand, no sign of anomaly appears in his second specimen, which was of rather higher purity than our specimen K 7.

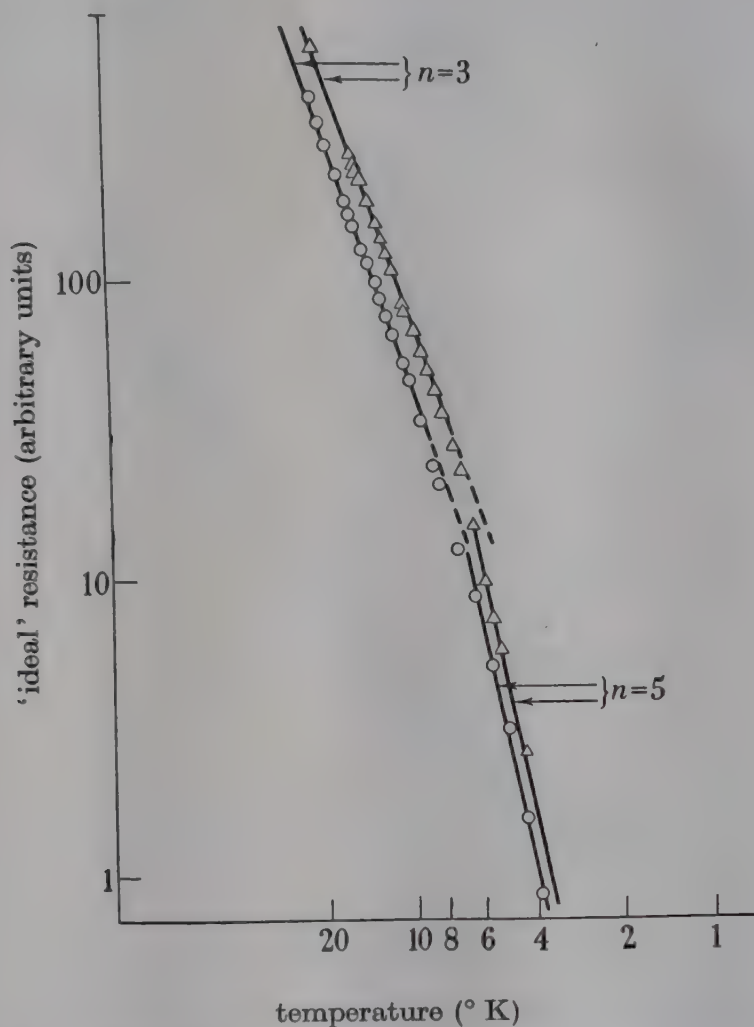


FIGURE 9. Resistance of potassium specimens plotted logarithmically.
Specimens: $\circ$, K 5; $\triangle$, K 7.

Our purer specimens (K 5 and K 7) are plotted logarithmically (in figure 9), and enable an accurate determination of power-law variation to be made down to $\sim 4^\circ$ K. The data of the purest (K 7), in particular, show an excellent T^3 variation down to 8° K and a T^5 law from 6.5° K downwards, as shown. This latter agrees well with van den Berg's observations of a T^5 law from 3 to 7° K, although he suggests an increase above the cubic power law below 14° K.

- * Chance Hysil G.H. 1: SiO_2 : 80.4% B_2O_3 : 12.4%,
 Na_2O : 4.2% CaO : 0.3%,
 Al_2O_3 : 2.7%.
- Chance (soft) G.W. 1: SiO_2 : 74% Na_2O : 16%,
 Al_2O_3 : 1% CaO : 7%,
 Sb_2O_3 : 1.3%.

(4) *Rubidium*

The metal was obtained from Messrs New Metals and Chemicals Ltd. This had been distilled to a high purity for use in photo-electric cells. The ratio of resistance in liquid helium at 4.2°K to that at room temperature was about 3×10^{-2} and 1.8×10^{-2} for our two specimens, which are rather better values than those of the specimens used by Meissner & Voigt.

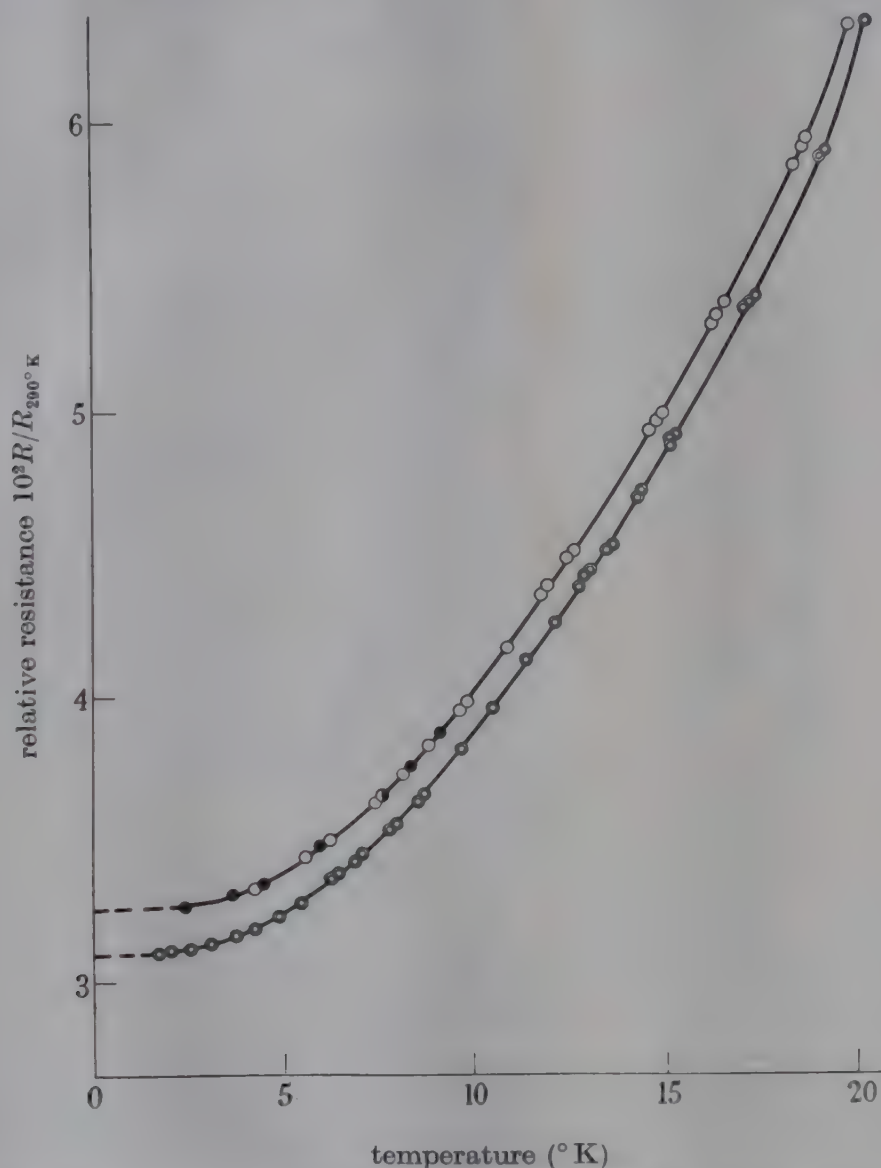


FIGURE 10. Resistance of rubidium (specimen Rb 1) below 20°K . Dates of experiment (1948):
 O, 22 January; ●, 29 January; ⊙, 5 February.

A number of experiments were made from 20°K downwards; no significant anomaly was detected anywhere in this region.* We may say at once that the resistance does not follow Grüneisen's law, but since on general grounds (cf. Seitz 1940, p. 532) one should expect a T^5 law to hold in any case for *sufficiently* low temperatures, a detailed examination was made in the liquid helium region down to $\sim 1.5^\circ\text{K}$.

Experimental results for rubidium are shown in figure 10, and logarithmic plots for both specimens are given in figure 11. Meissner & Voigt did not attempt to

* Apart from possibly a *very* slight extended 'hump' from 7 to 12°K .

compare their data with equation (1) but tried to fit an earlier empirical law due to Grüneisen, viz. $R \sim ATC_p$, where A is a constant and C_p the specific heat. This leads to $R \sim T^4$ for low temperatures. They were probably led to this course, since the decline of the resistance with temperature is much too slow for the Bloch-Grüneisen law with constant Θ . Even then the agreement is not good. We have therefore derived a curve of the variation with temperature of the required Debye temperature, Θ (see figure 7), in the same fashion as one exhibits specific heat data that do not follow the Debye law strictly. One sees immediately that the data for rubidium are highly anomalous. It should be pointed out that the reality of the

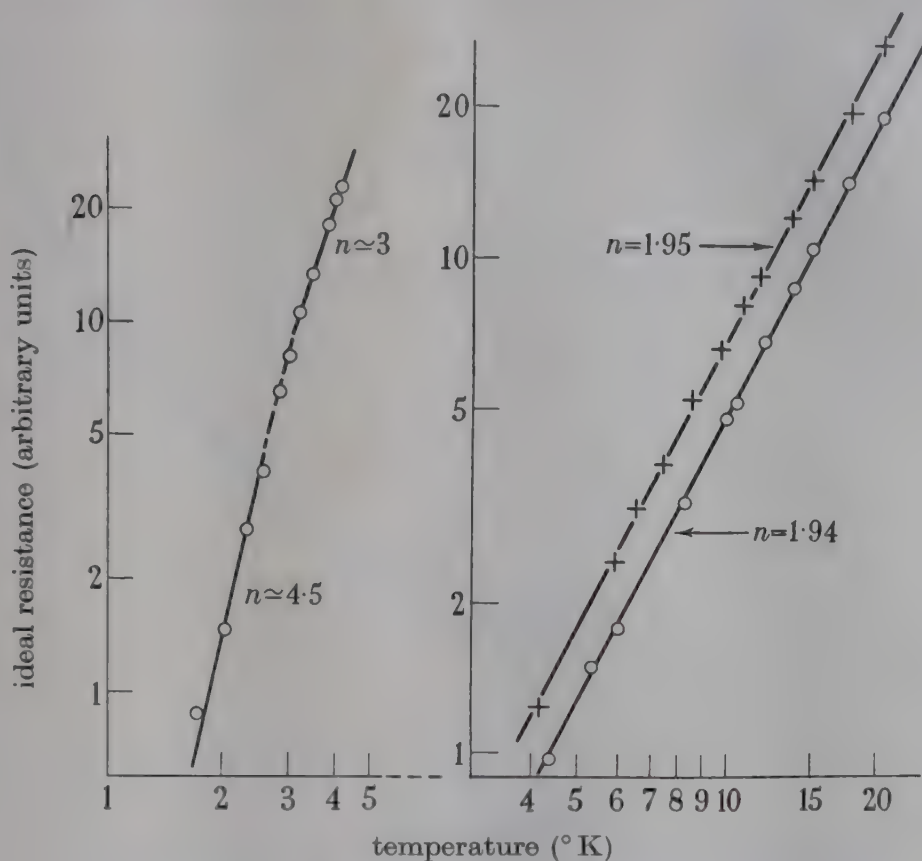


FIGURE 11. Logarithmic plot of resistance of rubidium.

minimum above 20°K cannot be guaranteed, since the value of Θ deduced in this region is found to be rather sensitive to temperature and high accuracy is not claimed for temperature measurement in this region. If the deviation from theory of the electrical resistance for the alkali metals, except sodium, is in fact ultimately to be attributed to the departure of the lattice vibrations from the Debye law, then this curve is of significance and suggests that a determination of the specific heat of rubidium would be of very considerable interest; the same applies with equal weight to caesium. In this respect the qualitatively similar trend of Θ for Li when deduced from both sources of data may be significant. On the other hand, if the source of the anomalous behaviour is an intrinsic deviation of the electronic structure of a metal from that of the ideal metal, then of course the curve is only to be regarded as a convenient mode of data presentation.

Two experiments were made to examine closely the resistance variation below $\sim 4^\circ \text{K}$; these agree well with one another and are shown in figure 12. The data are

plotted logarithmically in figure 11, and indicate clearly that at sufficiently low temperatures the T^5 law does hold. The fact, however, that this does not occur until probably below $\sim 2^\circ$ K indicates a very low limiting Debye temperature ($\sim 25^\circ$ K) in good agreement with the variation of Θ shown in figure 7. Of the various possible sources of anomalous behaviour in general, only electron-electron collisions (Baber 1937) leading to a T^2 component of resistivity appear likely to prevent the ultimate inception of a T^5 law, and such collisions would be expected to play a very small part in the alkali metals.

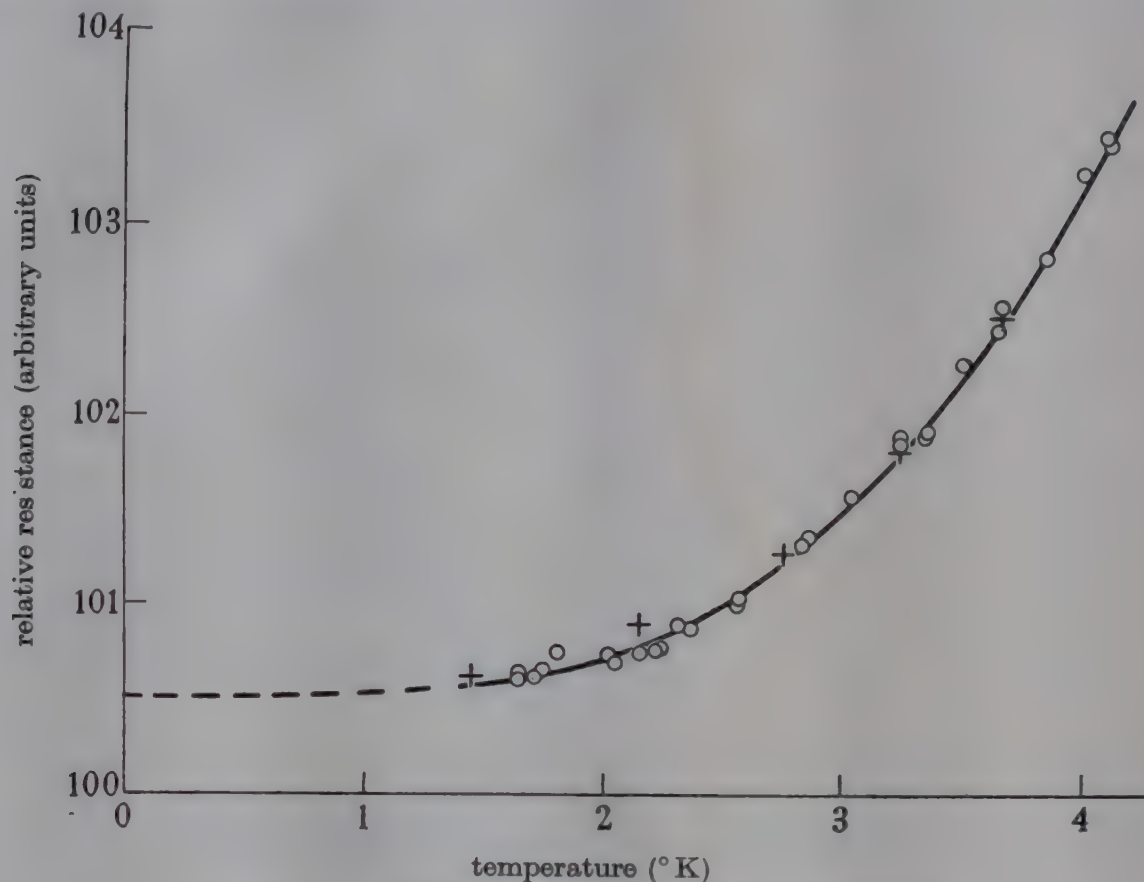


FIGURE 12. Resistance of rubidium at low temperatures.
+, Experiment 1; O, Experiment 2.

(5) Caesium

This metal was also obtained from Messrs New Metals and Chemicals Ltd., and the preparation before despatch to us was similar to that for rubidium. Two capsules containing about 1 g. each were received, and spectrographic analyses obtained through the courtesy of A.E.R.E. Harwell furnished the following data:

sample I (Cs 1)		sample II (Cs 2)	
	parts in 10^6		parts in 10^6
K	500	Na	200
Si	500	K	200
Rb	300	Rb	300
Pb	300	Pb	50
Na	1000	<i>Not detected</i>	
Li	5	Li	(< 1)
Sn	20	B	(< 50)
<i>Not detected</i>		Hg	(< 200)
B	(< 50)	Sn	(< 10)
Hg	(< 200)	Si	

Judging also from the residual resistance ratios (2 to 1×10^{-2}) the metal was rather purer than that used by Meissner & Voigt.

The first results on caesium showed quite anomalous behaviour. The curve of resistance from 20 to $\sim 1.5^\circ \text{K}$ showed two marked 'kinks' (cf. MacDonald & Mendelssohn 1948), the first at $\sim 6^\circ \text{K}$, the second, more severe, at $\sim 4^\circ \text{K}$. Subsequent experiments clearly confirmed these anomalies. A second specimen, from

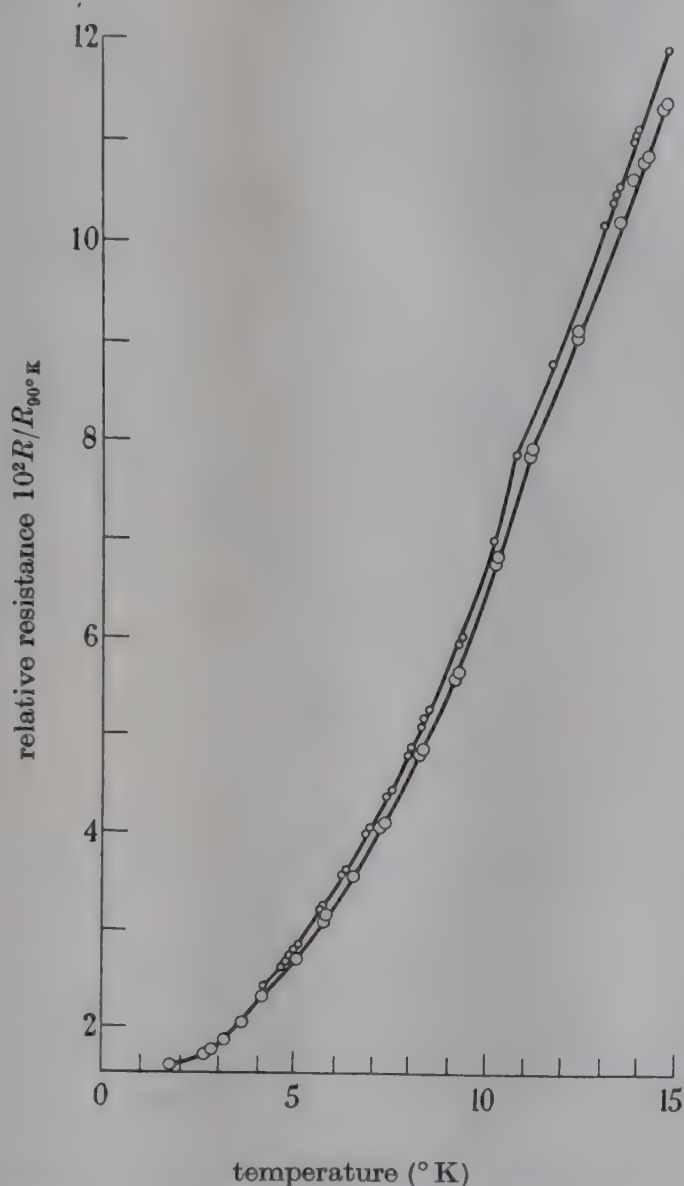


FIGURE 13. Resistance variation at low temperatures of caesium (specimen Cs2).
Dates of experiment (1948): $\circ$, 19 February; $\bigcirc$, 26 February.

the other capsule of metal, also showed anomalies, somewhat less marked, at ~ 4 and $\sim 11^\circ \text{K}$ (see figure 13). It is interesting to note that Meissner & Voigt suspected a resistance anomaly between 4 and 20°K .

When relatively slowly cooling the first specimen to liquid oxygen the resistance showed a discontinuity; the specimen was re-warmed to room temperature and then successfully carried down to liquid oxygen temperature by extremely slow cooling. After the low-temperature work, investigations were made between -183°C and room temperature using a thermo-couple for the temperature determination. One of these experiments is shown in figure 14, and a discontinuity on first warming is

seen at $\sim -20^\circ\text{C}$ which did not arise, after re-cooling, in the second experiment up to room temperature. It is clear that this anomaly is responsible for the radically different behaviour of the specimens observed by McLennan, Niven & Wilhelm (1928) and by Meissner & Voigt in this temperature region.

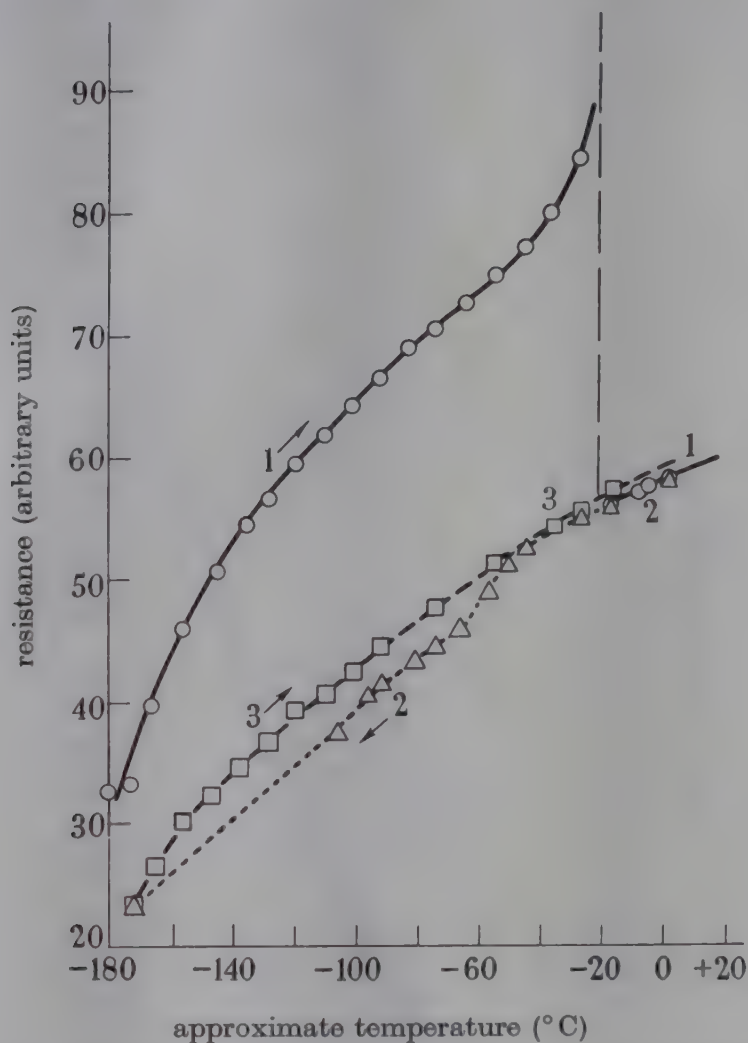


FIGURE 14. Variation of resistance of caesium (specimen Cs1) at temperatures above liquid oxygen. Temperature scale based on linear interpolation of Cu-constantan thermocouple:

experiment	key	temperature
1	○	increasing
2	△	decreasing
3	□	increasing

The electrical behaviour at -20°C may possibly be ascribed to some anomaly in the thermal contraction in this region, since it appears likely that only such a mechanism could produce effectively infinite resistance.

In view of the unusual behaviour of the metal, in particular at very low temperatures, leading to uncertainty in the true value of residual resistance it does not appear profitable to try to compare the data in detail with the theoretical conductivity law. We mention, however, that ignoring the anomalous behaviour, the apparent Debye temperature for Cs1 deduced from data over the range ~ 8 to 20°K is only 35°K (cf. figure 15), while even below 4°K the resistance follows an approximate power law of index only ~ 2.75 , obviously still deviating markedly from the Bloch law.

Since the sodium content of Cs 1 is considerably greater than in Cs 2 (see analysis) it seems likely that the low temperature anomalies are again attributable to this source as in the case of potassium.

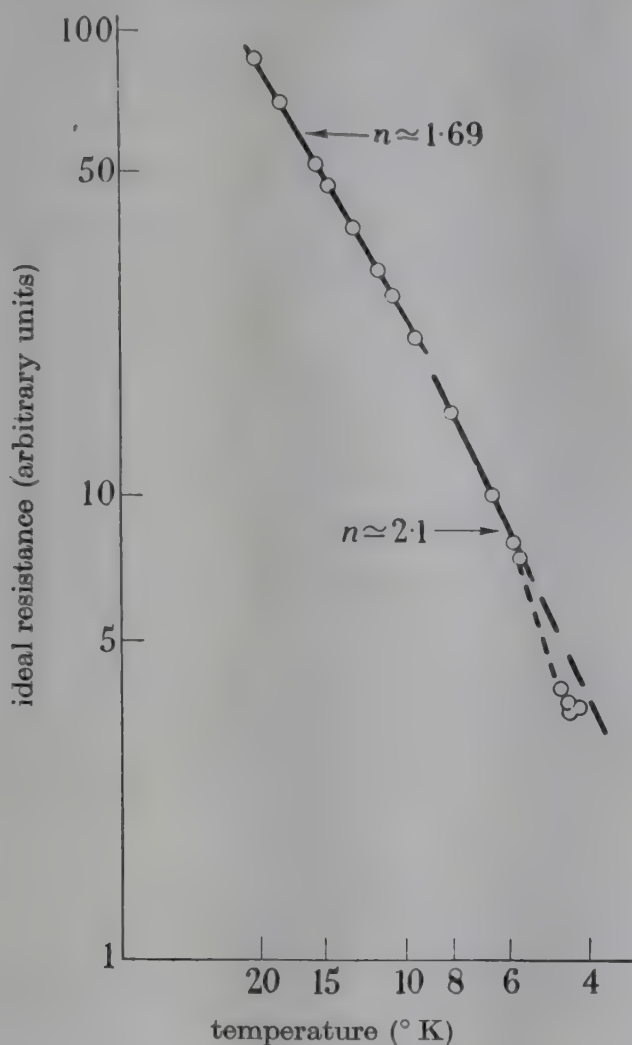


FIGURE 15. Logarithmic plot for caesium (specimen Cs 1).

CONCLUSIONS

The specific features of the results have already been discussed in the case of each metal in turn. Surveying the data generally under the headings of the introduction it appears that sodium follows the theory remarkably well throughout. This confirms indications from other sources that sodium may be regarded as an excellent example of the quasi-free electron model in all respects. The results on the other alkali metals seem to show, however, that sodium presents a singular case in this respect. There are two ways in which the deviations from the simple theory can be interpreted. On the one hand the lattice vibrations may depart markedly from the Debye law, resulting in a different law of variation of resistance with temperature. The second possibility is that the conduction electrons may not be legitimately regarded as quasi-free. In the case of lithium the screening of the valence electron from the nucleus may be inadequate. On the other hand, in the later alkalis the growth of the ion may be sufficient to interfere with the freedom of the outer electron.

An entirely new phenomenon is presented by the small anomalies observed in potassium and caesium which have been correlated with the presence of very slight impurity. Although the phenomena are not conspicuous they appear to be of particular interest because so far impurity scatter has not been known to create any effect except the residual resistance. This temperature-independent component of the resistance which forms the basis of Matthiessen's rule can under plausible assumptions be accounted for theoretically (cf. Mott & Jones 1936, pp. 287-288), but there does not seem to be any simple way of interpreting these small anomalies in a similar way. It is hoped to obtain further information on these questions from resistance measurements on other metals, and these are now in progress.

REFERENCES

- Baber, W. G. 1937 *Proc. Roy. Soc. A*, **158**, 383.
 Blackman, M. 1937*a* *Proc. Roy. Soc. A*, **159**, 416.
 Blackman, M. 1937*b* *Proc. Camb. Phil. Soc.* **23**, 93.
 Bloch, F. 1929 *Z. Phys.* **59**, 208.
 Bridgman, P. W. 1931 *The physics of high pressure*. London: International Text Books of Exact Science.
 van den Berg, G. J. 1938 Thesis, Leiden.
 Daunt, J. G. & Mendelssohn, K. 1938 *Proc. Phys. Soc.* **50**, 525.
 Frankenhaeuser, B. & MacDonald, D. K. C. 1949 *J. Sci. Instrum.* **26**, 145.
 Fuchs, K. 1936 *Proc. Roy. Soc. A*, **153**, 622.
 de Haas, W. J. & van den Berg, G. J. 1936 *Commun. Phys. Lab. Univ. Leiden*, Suppl. no. 82*a*.
 Kapitza, P. & Milner, C. J. 1937 *J. Sci. Instrum.* **14**, 165.
 Kellerman, E. W. 1940 *Phil. Trans. A*, **238**, 513.
 Kellerman, E. W. 1941 *Proc. Roy. Soc. A*, **178**, 17.
 MacDonald, D. K. C. 1947 *J. Sci. Instrum.* **24**, 232.
 MacDonald, D. K. C. 1950 *J. Sci. Instrum.* (In preparation.)
 MacDonald, D. K. C. & Mendelssohn, K. 1948 *Nature*, **161**, 972.
 McLennan, J. C., Niven, C. D. & Wilhelm, J. O. 1928 *Phil. Mag.* **6**, 672.
 Meissner, W. 1920 *Phys. Z.* **2**, 373.
 Meissner, W. & Voigt, B. 1930 *Ann. Phys, Lpz.*, **7**, 761, 892.
 Mendelssohn, K. 1931 *Z. Phys.* **73**, 482.
 Mendelssohn, K. 1936 *Proc. Roy. Soc. A*, **155**, 558.
 Milner, C. J. 1937 *Proc. Roy. Soc. A*, **160**, 207.
 Mott, N. F. & Jones, H. 1936 *Theory of the properties of metals and alloys*. Oxford University Press.
 Preston, J. S. 1946 *J. Sci. Instrum.* **23**, 173.
 Schottky, W. 1922 *Phys. Z.* **23**, 448.
 Seitz, F. 1940 *The modern theory of solids*. New York: McGraw-Hill.
 Simon, F. 1926 *S.B. preuss. Akad. Wiss.* p. 447.
 Simon, F. 1930 *Ergebn. exakt. Naturw.* **9**, 223.
 Simon, F. & Pickard, G. L. 1948 *Proc. Phys. Soc.* **61**, 1.
 Woltjer, H. R. & Kamerlingh Onnes, H. 1924 *Commun. Phys. Lab. Univ. Leiden*. No. 173*a*.

The fine structure of the hydrogen α line

By H. KUHN AND G. W. SERIES, *Clarendon Laboratory, University of Oxford*

(Communicated by Lord Cherwell, F.R.S.—Received 23 January 1950)

[Plate 12]

A new attempt has been made to resolve spectroscopically the fine structure of the α line, 6563 Å, of the Balmer series of heavy hydrogen. A deuterium discharge tube, whose walls were partly of copper, was immersed in liquid hydrogen and run at very low current densities. By this means the Doppler width of the individual lines was sufficiently reduced to enable the satellite between the main components of the 'doublet' to be resolved much more completely and to be measured more accurately than was possible before. The spectroscopic resolution was achieved with a Fabry-Perot interferometer of large diameter, mounted internally in a prism spectrograph.

The measurements support the assumption that the $2S_{\frac{1}{2}}$ term is higher than is predicted by the Dirac theory, and lead to the value $0.0369 \pm 0.0016 \text{ cm.}^{-1}$ for the amount of the shift, in good agreement with the latest value found by Lamb & Retherford from microwave measurements. The relative widths of the main components and the interval between them also provide qualitative evidence for the shift of the $2S_{\frac{1}{2}}$ term. The intensities of the lines were found to agree approximately with the prediction of the Dirac theory.

The weak line $2P_{\frac{3}{2}}-3S_{\frac{1}{2}}$ was partly resolved from the neighbouring strong line $2P_{\frac{3}{2}}-3D_{\frac{3}{2}}$. The provisional value for the distance between them would indicate that the $3S_{\frac{1}{2}}$ term is shifted very little, if at all. Further attempts are being made to obtain a more reliable value for this term.

1. INTRODUCTION

The structure of the first line of the Balmer series, 6563 Å, has been the object of a large number of investigations and controversies. Most workers confined their attention to the measurement of the wave-number difference of the two main peaks, generally called (1) and (2), comparing it with the prediction of current theories. Though the existence of a third component, (3), between the main peaks, was deduced as early as 1925 by Hansen from an analysis of his photometer tracings, the large Doppler width, due to the small mass of the hydrogen atom, prevented the resolution and measurement of this third component.

The existence of this component was subsequently explained by the introduction of the electron spin into the theory. The combined effect of the spin and relativity corrections in the wave-mechanical theory of the hydrogen atom (Goudsmit & Uhlenbeck 1925; Heisenberg & Jordan 1926; Pauli 1927) led to the same values of the terms as Sommerfeld's theory (1916), but altered their interpretation and the selection rules. Dirac (1928) derived the same results in a more consistent way, and the term 'Dirac theory' will be used in this paper when these results are referred to.

In heavy hydrogen, ^2H , the fine structure is expected, on theoretical grounds, to be essentially the same as in light hydrogen, but the Doppler width is 1.4 times smaller, and the use of this isotope for investigations of the fine structure made considerable advances possible. The best resolution was obtained by R. C. Williams (1938*a*), whose photometer tracings show a very slight dip between the components (3) and (2). This partial resolution enabled him to make an approximate measurement of the position of the component (3) which, he concluded, was not in the

position predicted by the Dirac theory. The spacing between the main peaks (1) and (2) also disagreed with the prediction of this theory. Pasternack (1938) suggested that these discrepancies could be explained by assuming an upward displacement of the $2S_{\frac{1}{2}}$ level of about 0.03 cm.^{-1} from its theoretical position. The work of other observers (Drinkwater, Richardson & Williams 1940), however, contradicted these results, which were therefore not generally accepted until the shift of the $2S_{\frac{1}{2}}$ level was confirmed by Lamb & Retherford (1947) using a microwave method based on the metastability of the $2S_{\frac{1}{2}}$ state. These experiments were carried out with light hydrogen, but deuterium was later found to give the same results.

In most of the previous spectroscopic work, discharge tubes immersed in liquid air had been used. Only one attempt to use liquid hydrogen for cooling a discharge tube appears to have been reported (Heyden 1937), but the resolution was not improved. The application of the atomic beam technique to atomic hydrogen has been discussed and tried (Mack & Barkowsky 1942; Williams 1942), but apparently without success.

In the present investigation of the fine structure of the α line of deuterium, a specially designed discharge tube with walls partly of copper was immersed in a bath of liquid hydrogen. With the extremely small current densities used, it was found that temperatures of the emitting atoms below 55° K were reached, resulting in a considerable improvement in resolution. It was thus possible to measure the position of the third component directly, and it was also found that the component (2) was wider than (1), owing to unresolved structure. Our results prove the reality of the shift of the $2S_{\frac{1}{2}}$ level and give the amount of this shift to an accuracy of about 5%.

Previous workers had often found large discrepancies between the observed intensity ratios of the components and the predictions of the theory. In the conditions of the experiments described below, the intensities differ only little from the theoretical values. This fact removes some objections which could be raised against previous work, and indicates that the spectrum observed can be ascribed to unperturbed hydrogen atoms.

The faint component $2P_{\frac{3}{2}}-3S_{\frac{1}{2}}$ was resolved and its position measured, though not with the same accuracy as for the other components.

2. THE LIGHT SOURCE

The design of the discharge tube which was used in the final experiments is shown in figure 1. It consists partly of copper and partly of Pyrex glass, using commercial copper-glass seals. A mixture of helium and heavy hydrogen passed vertically downward through the copper tube *A*, where it was precooled before entering the zone *C* of the cathode glow. It then passed through the short glass tube *G* to the copper section *B*, at the entrance to which a positive column *P* was formed. The light from the latter emerged through the lens *L* and the window *W*. The positive column was only 2 to 3 cm. long. Complications through self-absorption, which often arise in end-on observations, were thus avoided.

The design and operation of the discharge tube were based on the following considerations:

(1) It is clear that the use of copper walls instead of glass walls ensures a considerably greater rate of heat flow into the cooling bath and thus helps to reduce the temperature of the gas. The tendency of metal walls to favour recombination of atoms can be overcome by covering the walls with a thin layer of heavy ice (Wood 1921).

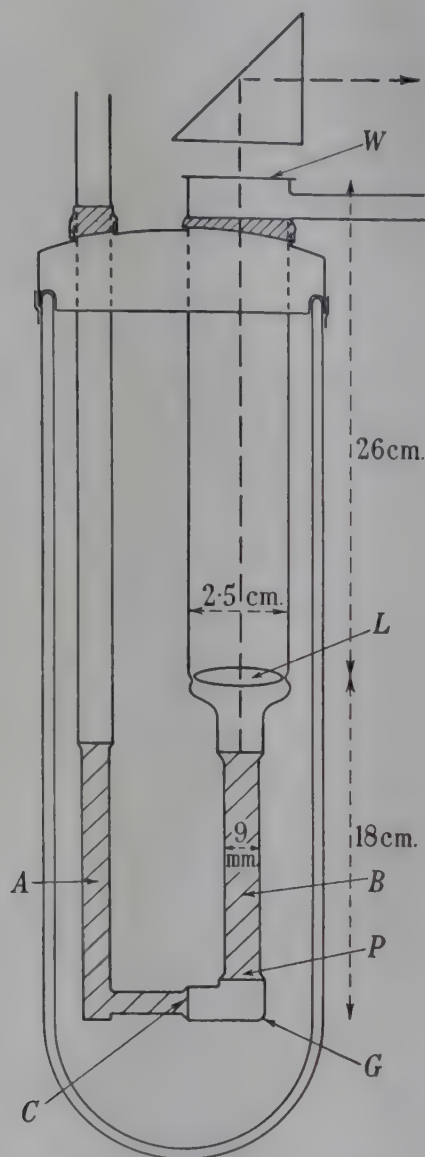


FIGURE 1. Discharge tube.

(2) The hydrogen atoms are formed from molecules by excitation and simultaneous dissociation, and have initially a large amount of kinetic energy. If the density of helium is high enough, collisions with helium atoms will rapidly reduce the temperature of the hydrogen atoms (the mean free path was about $\frac{1}{100}$ mm. under the normal conditions of the discharge). Most of the hydrogen atoms which are available for excitation will therefore have low temperature. The process of excitation of the atom by electron impact will also impart some kinetic energy to the atom; but an estimate shows that the effect is small and would only become important at temperatures below those reached in the present experiments.

(3) The formation of atoms will be strong in the cathode glow. The gas, when reaching the anode, can thus be expected to contain a certain amount of deuterium in the atomic state. It is difficult to say how far this process of preformation of atoms was actually taking place under the conditions of our experiments.

(4) Wide aperture of the emergent beam of light was essential, in view of the wide aperture of the spectrograph used (see § 3). In order to reconcile this requirement with the necessarily large distance from the positive column to the window, the lens L was inserted inside the discharge tube.

Before the tube was filled, heavy-water vapour was admitted. When the cooling bath was then applied, a film of heavy ice formed on the walls.

The most favourable values of the pressures of helium and deuterium were found by trial. Too high pressure of deuterium made the molecular spectrum of deuterium appear too strongly, and too high pressure of helium produced He_2 bands. In most experiments, the pressure of helium was about 1.7, that of deuterium 0.2 mm. Hg, though other pressures were also used. The pressure was measured by means of a sensitive glass gauge of the Bourdon type ('spoon gauge'). A mercury diffusion pump maintained the circulation of the gas mixture, which, before entering the discharge tube, passed through a trap with charcoal, cooled with liquid air.

Figure 1 shows how the discharge tube was fitted into a brass cap which formed the lid of the Dewar vessel containing the liquid hydrogen. The anode limb consisted of a 2.5 cm. wide glass tube, the upper end of which protruded through the metal cap and was closed by a plane window. The clearance between the lens L and the walls of the glass tube was enough to allow the flow of gas towards the pump.

A connexion from the brass cap to a rotary oil pump allowed the hydrogen bath to be used at low pressure. The temperature of the bath was reduced to 14° K in some experiments.

The voltage applied to the discharge tube in normal running conditions was always below 300 V; the current was varied from 2 to 20 mA, but the 'standard' current used in the exposures on which the final results were based was 5 mA. The current density was then about 7 mA/cm.².

In the earlier stages of the work, a discharge tube of different design was used; the cathode was situated above the anode in the wide, vertical glass tube. It consisted of an aluminium ring almost touching the wall of the tube. The light from the positive column passed through the cathode ring; since an image of the positive column was formed on the slit, no light from the negative glow surrounding the ring was expected to reach the slit. But this method was open to objections and was abandoned in favour of the design described above.

3. THE OPTICAL SYSTEM

Liquid hydrogen is a more effective cooling agent than liquid air only if the power input into the discharge tube is exceedingly small. The optical system therefore had to be designed so as to produce the greatest possible intensity of the image for the given, very small, brightness of the source. At the same time, the resolving power had to be great enough to avoid any appreciable loss of the resolution which the

reduced Doppler width allowed. But the resolving power could not be increased by increasing the spacing of the etalon plates, since this would have entailed the overlapping of adjacent orders. A Fabry-Perot interferometer of very large diameter was therefore used, and the silver films were made just thick enough to produce the required, previously calculated, reflectivity. It was combined with a prism spectrograph constructed in the laboratory workshop and designed specially for this purpose.

The focal length of the camera lens of the spectrograph was chosen just large enough to separate, with the available prism, the H_α line from certain adjacent deuterium bands and from the helium red line 6678 Å, with a width of the slit image of 1 mm. A telescope lens of focal length 65 cm. satisfied this condition, in conjunction with a 60° prism, a combination which produced a dispersion of 100 Å/mm. in the red. The height of the prism was 9 cm.

A focal length of 65 cm. is sufficient to produce, with an etalon of spacing 7.67 mm., fringes of a large enough scale for photometric work with plates of medium grain size. Since, further, the size of the available etalon was similar to that of the prism, internal mounting was chosen. The etalon was placed between prism and camera lens. The collimator was a telescope lens of focal length 44 cm.

The image of the source formed on the photographic plate had to be large enough to avoid serious variation of intensity within the fringe pattern. This required a linear magnification of about 3:2. The effective aperture of the beam leaving the source had therefore to be at least $\frac{3}{2}$ times larger than the effective aperture of the camera lens which was about 1:10. By the device of placing a lens inside the discharge tube, an aperture of 1:6.5 was obtained.

The Fabry-Perot etalon consisted of two plates of fused silica, of diameter 11.5 cm., made by Messrs Adam Hilger Ltd. The spacers were also made of fused silica. An area of diameter 7 cm. was generally used, with the exception of a small portion which had to be masked off. It was found, however, that further reduction of the exposed area improved the resolution slightly.

The plates were silvered by evaporation in a vacuum and had a reflectivity of 91 % and a transmission factor of 5 % for red light. This reflectivity gives a theoretical half-value width of $\frac{1}{33}$ of an order. With the spacer of 7.67 mm. which was used in most experiments, this corresponds to 0.02 cm.⁻¹. Test exposures with cadmium red light showed that the actual resolution was not much below this value.

The spectrograph and interferometer were placed in a room, the temperature of which was thermostatically controlled to within $\frac{1}{5}$ to $\frac{1}{10}^\circ$ C. The discharge tube was set up in the adjacent room; the light entered through a slot in the wall containing a small glass window. Under favourable conditions, this arrangement permitted exposures of several hours, though generally the time of exposure was only 15 to 30 min.

Ilford H.P. 3 plates were mostly used. They proved to be faster, for light of wave-length 6563 Å, than other types of plates tried (Astra VIII, Kodak IHA, Astra VII). Of these, only Astra VII had considerably finer grain, but it was found to require too long exposures.

4. THE MEASUREMENT OF THE WAVE-NUMBER DIFFERENCES

Figure 2, plate 12, gives a typical spectrogram taken with the 7.67 mm. spacer. It shows the components (1), (2) and (3) in three different orders.

The main object of this work was the accurate measurement of the spacing between the components (1) and (3), of which one is six times stronger than the other. When the exposure was long enough to make the fainter line clearly visible, the stronger one was over-exposed, even on the comparatively soft plates used (Ilford H.P. 3). Only in a correctly exposed fringe can the setting of the cross-wire in the comparator be expected to give the position of the maximum; in an over-exposed fringe, any slight asymmetry, due to the genuine shape of the line or to instrumental effects, will give rise to systematic differences between the setting and the true maximum. Direct measurement of the distance between the two lines photographed in one single exposure would thus have been liable to serious errors. The following technique was therefore developed: three successive exposures were made on one plate, with times of exposure in the ratio 1:10:1. The distances of the strong fringes (1) and (2) from a reference mark were measured in the first and last exposure, and the distances of the faint fringes, due to component (3), from the same reference mark in the middle exposure. The plate was rejected if the distances measured in the first and third exposure did not agree well enough.

After some preliminary experiments with a glass fibre across the slit, suitable reference marks were ultimately found in the outer fringes which are narrower than those used for the measurement of the fine structure and could therefore be measured very accurately. In the second, longer exposure, the end of the slit which framed the reference fringes was covered during most of the time, and was only uncovered for 1 or 2 min. at the beginning and at the end of the exposure.

The centre of the fringe system can be found directly if fringes on both sides of the centre are visible. Such a setting would, however, have resulted in a considerable sacrifice in the number of orders in which the fine structure could be measured, and fringes on one side only of the centre were usually used. The centre was calculated, in the customary way, from the measured position of the fringes in three orders. This method relies, to some extent, on the validity of the square-root formula for the radii of the fringes, and thus on the absence of distortion in the camera lens. A series of measurements of interference fringes produced with cadmium red light showed that the formula held, in fact, accurately.

Of the two main peaks of the fine structure, that of shorter wave-length, (2), is a blend of two lines of similar intensities (see § 6), and thus not so suitable for measurement and for an analysis of the structure. The wave-number difference of components (3) and (1) was therefore chosen as the principal aim of the measurements. The wave-number difference (2) - (1) was also measured, partly for comparison with results of other workers.

In an attempt to study the possible influence of variations of discharge conditions, especially of the current density and deuterium pressure, a large number of exposures were taken in which these conditions were varied over a wide range.

With increase in the current, the resolution of component (3) decreased, and the

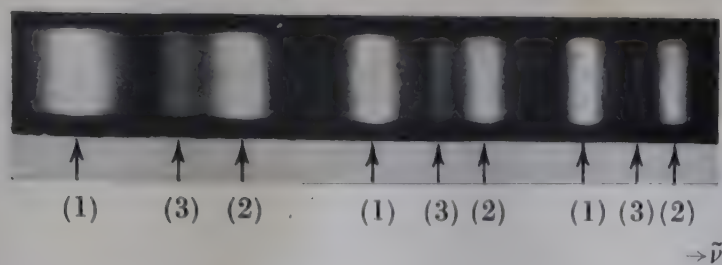


FIGURE 2. Fine structure in three orders, with 7.67 mm. spacer; current 5 mA.

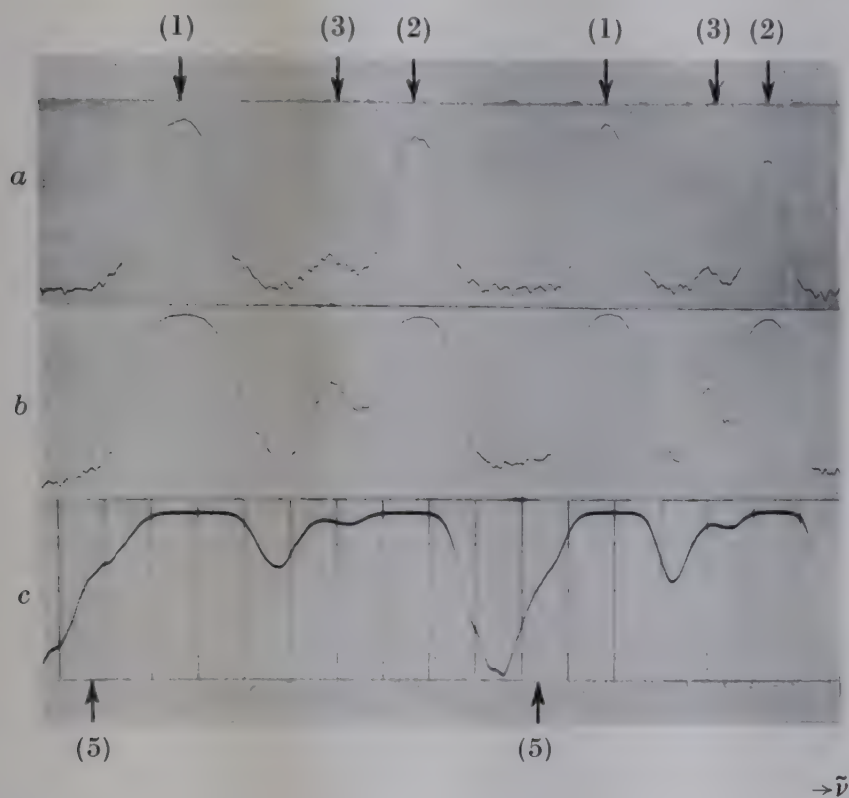


FIGURE 3. Photometer tracings, 7.67 mm. spacer; current 5 mA.
a, short exposure; b, standard exposure; c, long exposure.

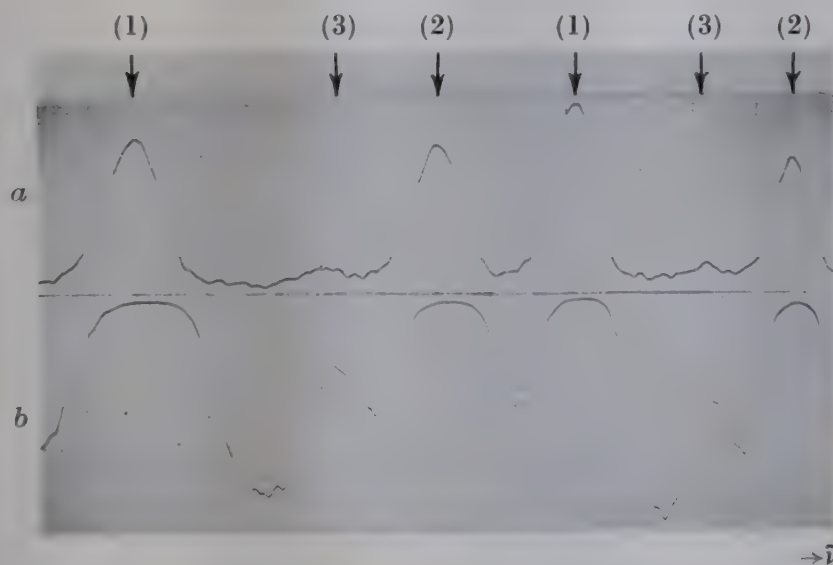


FIGURE 4. Photometer tracings, 10 mm. spacer, cooling bath at 14° K.
a, current 2 mA; b, current 5 mA.

measurements became less accurate, but no change in the distances could be observed, as shown in table 1. The number of fringes from which each value has been derived is given in brackets.

In order to test the influence of the variation of the pressure of deuterium, three sets of exposures were taken. In the first, the deuterium pressure was so low that the deuterium α line appeared weaker than the helium red line, in the second set the deuterium pressure had the standard value of 0.2 mm., whereas in the third it was so high that the helium red line had completely disappeared. The three average values, each taken over six fringes, of the wave-number difference (3)–(1) were 0.188, 0.186 and 0.188 cm^{-1} , and for the difference (2)–(1), 0.325, 0.324 and 0.323 cm^{-1} . From these and a number of similar comparisons, it was concluded that a variation of the concentration of deuterium, sufficient to alter considerably the conditions of the discharge, did not affect the position of the components.

TABLE 1. SPACINGS OF COMPONENTS AT DIFFERENT CURRENTS

current (mA)	2–2½	5	10	20
(3)–(1)	0.188 (5)	0.189 (16)	0.189 (6)	0.188 (15) cm^{-1}
(2)–(1)	0.323 (5)	0.322 (16)	0.323 (12)	0.322 (17) cm^{-1}

Comparisons between exposures in which the temperature of the cooling bath was 20 and 14° K respectively, i.e. made with liquid hydrogen at atmospheric and at reduced pressure, also failed to show any significant differences in the measured spacings. At the lower temperature, the lines were noticeably narrower when the current was below 5 mA.

Great attention was paid to the possible influence of traces of light hydrogen which are impossible to exclude in an apparatus containing greased taps. Occasional tests with an etalon of suitably chosen spacing showed that the proportion of light to heavy hydrogen was less than 1:50. With the spacing of 7.67 mm., the component (1) of ^1H falls well outside the pattern of ^2H , and component (2) of ^1H midway between (1) and (3) of ^2H , well resolved from either. Even a considerably greater amount of light hydrogen than that actually present would therefore not have affected the measurements. The mentioned relative positions of the lines were calculated from the known isotope shift (Williams 1938*b*; Drinkwater *et al.* 1940) and verified in an exposure taken after a splinter of sealing wax had fallen into the discharge tube, as the result of an accident; the lines of ^1H appeared then in the calculated positions.

Though variations in discharge conditions had no measurable effect on the positions of the lines, the most reliable results were clearly those taken at very low current density, since they give the best resolution and consequently the greatest accuracy of measurement, and are also least likely to be affected by electric fields and other possible causes of error (see § 10). The final values for the positions of the lines were therefore based on a number of exposures taken under 'standard conditions', with a current of 5 mA and with the cooling bath at atmospheric pressure. They are summarized in table 2, where each line represents one plate.

The three orders measured on each plate were produced by light coming from different parts of the cross-section of the discharge tube. No significant difference was found between the averages of the interval (3)–(1) measured in the different orders.

In order to determine accurately the spacing between the etalon plates, the interference fringes produced by the two mercury yellow lines were measured. A mercury lamp containing the pure isotope 198, kindly supplied by the General Electric Company, was used for this purpose. The spacing (multiplied by the refractive index of air) was found as 0.7675 cm.

TABLE 2. WAVE-NUMBER DIFFERENCES OF COMPONENTS IN CM.⁻¹

order counted from centre	...	(3)–(1)			(2)–(1)		
		2	3	4	2	3	4
current 5		0.1864	0.1868	—	0.3210	0.3199	—
mA. 5		0.1879	0.1875	0.1898	—	0.3204	0.3211
5		0.1876	0.1862	—	0.3215	0.3215	—
2.5		0.1882	0.1844	—	0.3196	0.3205	—
5		0.1865	0.1880	0.1867	0.3202	0.3206	0.3221
5		0.1889	0.1923	0.1859	0.3214	0.3226	0.3218
5		0.1881	0.1895	0.1882	0.3217	0.3207	0.3223
5		0.1904	0.1897	0.1888	0.3206	0.3224	0.3206
5		0.1861	0.1878	0.1858	0.3213	0.3228	0.3213
5		0.1879	0.1889	0.1875	0.3195	0.3232	0.3229
5		—	0.1868	—	0.3183	0.3232	0.3196
5		0.1875	0.1874	0.1874	0.3203	0.3221	0.3219
5		0.1857	0.1871	0.1866	0.3195	0.3222	0.3215
5		0.1895	0.1898	—	0.3209	0.3220	0.3236
5		0.1892	0.1881	0.1923	0.3210	0.3218	0.3210
mean		0.1879	0.1880	0.1879	0.3205	0.3217	0.3216
mean of all orders			0.1879			0.3213	
standard deviation			0.0003			0.0002	

Exposures were also taken with a 10 mm. spacer and gave good agreement with the values of table 2; but the component (2) of light hydrogen fell near the component (3) of deuterium, and very slight errors arising from this overlapping could not be excluded. The measurements were therefore not included in the final average.

In preliminary work, the wave-number difference (2)–(1) was also measured with a 5 mm. spacer, with results in close agreement with those in table 2.

The highest resolution was obtained with a current of 2 mA and a temperature of the bath of 14° K. It was especially in these exposures that the line (2) appeared noticeably wider than (1). From a tracing such as that shown in figure 4*a*, plate 12, the ratio of the widths was found to be about 1.3. The time of exposure was, however, very long with this low current, and component (3) was therefore not measured under these conditions.

Figure 3, plate 12, shows photometer tracings of spectrograms taken with the 7.67 mm. spacer, at a discharge current of 5 mA. The short exposure (*a*) shows qualitatively the relative intensities of components (1), (2) and (3). In the medium

exposure (b) the component (3) is seen more clearly. It represents the 'standard' conditions of current and exposure as used in most measurements. In the long exposure (c) the components (1), (2) and (3) are over-exposed, and component (5) appears as a hump on the wing of component (1).

Exposures with the 10 mm. spacer, and with the cooling bath at a temperature of 14°K , i.e. with the liquid hydrogen under reduced pressure, are shown in figure 4. The discharge currents were 2 and 5 mA in (a) and (b) respectively.

5. MEASUREMENT OF THE INTENSITIES

The spectral sensitivity curve of the Ilford H.P. 3 plate drops very rapidly at the wave-length of H_α , so that the usual method of calibration by means of white light in conjunction with a spectroscope was not accurate enough. The intensity marks were therefore made with light from a hydrogen discharge tube, from which the red line was selected by a colour filter. Neutral filters, previously calibrated, were used for producing known ratios of intensity. In this way, intensity marks were made on several of the plates on which the fine structure was photographed. The densities were measured both on the recording micro-photometer of the Oxford Observatory and on a non-recording Zeiss micro-photometer.

Line 3 of table 3 gives the measured ratios of the peak intensities of the components (1), (2) and (3). The intensity (1) was chosen as unit.

TABLE 3. RELATIVE INTENSITIES OF THE COMPONENTS

	(1)	(2)	(3)
theoretical, uncorrected	1.0	0.71	0.11
theoretical, corrected	1.0	0.65	0.12
experimental	1.0	0.70	0.18

6. THE POSITION OF THE $2S_{\frac{1}{2}}$ TERM

Relativistic quantum mechanics, either in the form of Schrödinger's theory with the inclusion of relativistic and spin corrections, or in the form of Dirac's theory, leads to the following expression for the term values of hydrogen:

$$T_{n,j} = -\frac{R}{n^2} - \frac{R\alpha^2}{n^3} \left(\frac{1}{j + \frac{1}{2}} - \frac{3}{4n} \right).$$

Figure 5a shows the terms with $n = 2$ and $n = 3$ plotted according to this formula, with the use of the latest numerical values of the constants $R_D = 109707\text{ cm.}^{-1}$ and $\alpha^2 = 5.3258 \times 10^{-5}$ (Du Mond & Cohen 1949). The intensities of the lines as calculated by Sommerfeld & Unsöld (1926) and Kupper (1928) are given as numbers next to the lines and are roughly indicated by the length of the lines.

Neglecting at first the small effect of the unresolved, weak component (4) and any instrumental corrections, it is clear that the measured distance (3)–(1) is about 0.03 to 0.04 cm.^{-1} smaller, and the distance (2)–(1) is about 0.007 cm.^{-1} smaller than the respective theoretical values. The only simple way of explaining these facts is the assumption that the $2S_{\frac{1}{2}}$ level is about 0.03 to 0.04 cm.^{-1} above the

$2P_{\frac{1}{2}}$ level, as shown in figure 5*b*. Pasternack (1938) had, in fact, made this assumption on the basis of the results of R. C. Williams (1938*a*). The new fact that component (2) appears wider than (1) on our spectrograms proves independently the qualitative correctness of this assumption. Accepting the fact that $2P_{\frac{1}{2}} - 2S_{\frac{1}{2}} = x$ is of the assumed order of magnitude, we can now derive an accurate value of x from an analysis of our measurements.

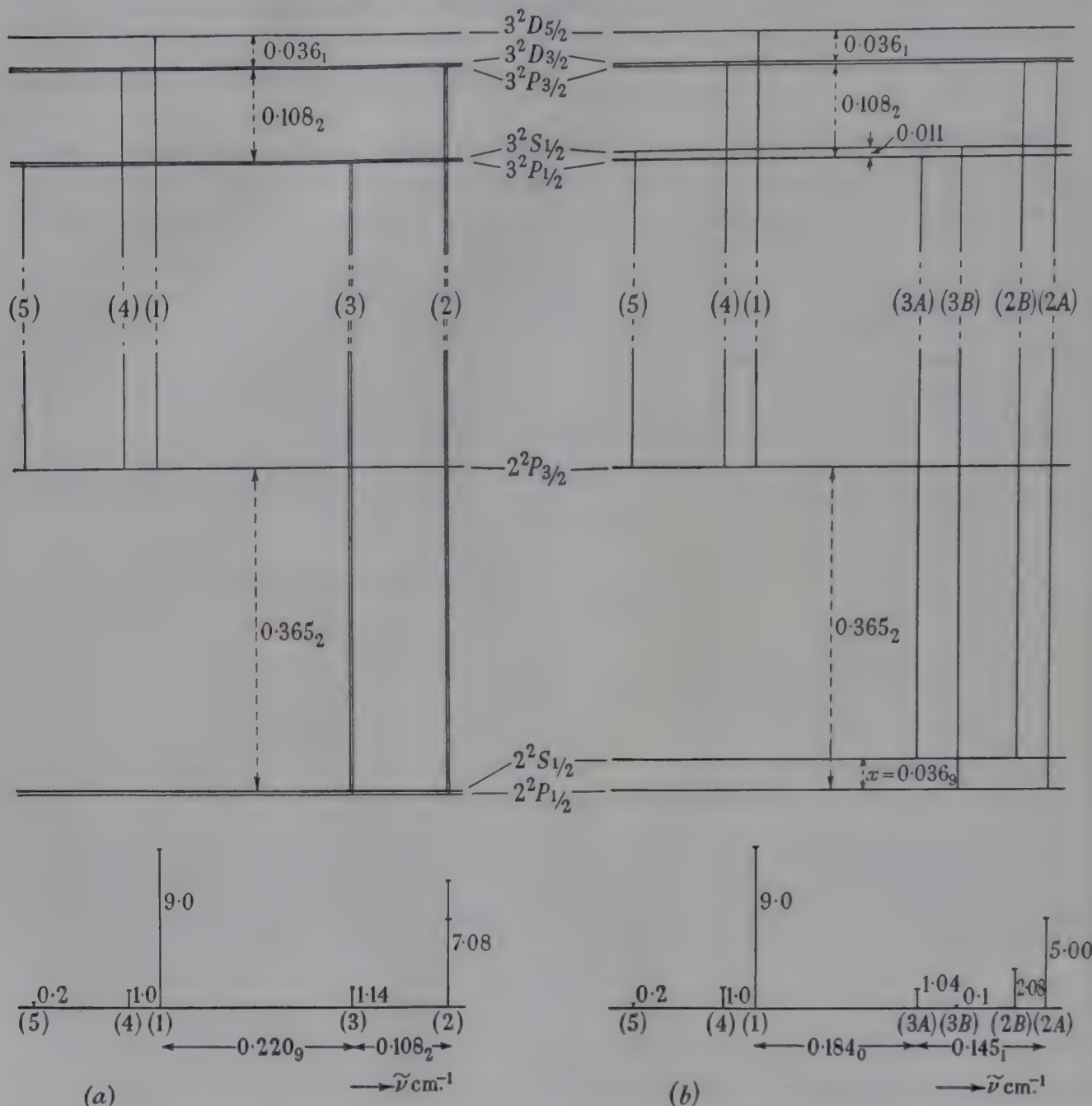


FIGURE 5. Term diagram. *a*, Dirac theory; *b*, showing shifts of S terms.

The intensities of the components in figure 5*b* are derived from the same theory as those in figure 5*a*. There is little doubt that these theoretical results for the intensities apply, even though the term values differ slightly from those given by the theory, if the cause of the difference can be ascribed to an internal effect in the atom, acting as a perturbation having central symmetry. The components (2) and (3) are now split into two subcomponents A and B. Though these splittings are not

resolved, the ratio of the intensities of the blends (2), (3) and (1+4) can be compared with the theoretical values. The first line in table 3 gives the theoretical intensities calculated by adding the subcomponents *A* and *B*. The figures in the second line are based on the same values for the individual components, but the peak intensity of unresolved components has been calculated by superposition of Doppler distribution curves. Some allowance has also been made for instrumental background, with the use of Airy's formula. But this estimate neglects the Doppler width of the strong lines and is therefore certainly too small. On account of this, and other possible sources of error, the values for (2) and especially for (3) in the second line of table 3 are probably rather too low. Comparison with the experimental values in line 3 does not indicate, therefore, any discrepancy outside the possible, rather wide, limits of error. In view of this, it can be assumed that the ratios of the intensities of the unresolved components also agree, at least roughly, with the theoretical values. In the experiments of R. C. Williams (1938*a*), the intensities differed very considerably from the theoretical values.

In figure 5*b*, the position of the term $3S_{\frac{1}{2}}$ has been assumed somewhat arbitrarily. This term leads to the faint component (3) *B* and therefore only affects a small correction applied to the position of component (3) (see below). The term difference $3P_{\frac{3}{2}}-3S_{\frac{1}{2}}$ has been assumed to be about one-third of the difference $2P_{\frac{3}{2}}-2S_{\frac{1}{2}}$. This may be justified either by reference to the new theories of radiation shifts of terms or by the general argument that the whole structure of the level $n = 3$ is about one-third as wide as that of the level $n = 2$.

In comparing the measured wave-number difference of the fringes (3) and (1) with the wave-number difference of the two single lines (3) *A* and (1), two effects have to be considered. (*a*) The peak of (1) will be slightly shifted by the presence of the unresolved, weak satellite (4) and the peak (3) *A* by the unresolved, weak satellite (3) *B*. (*b*) The lines (1) and (2), though well resolved from (3), are considerably stronger than the latter and therefore have a slight influence on its position. Both effects are small, particularly (*b*), and can therefore be estimated by comparatively crude methods.

The only important causes of the finite widths of the fringes are Doppler effect and instrumental broadening. The half-value width of the latter is considerably smaller than the half-value width of the Doppler effect, and the intensity distribution near the maximum will thus differ very little from pure Doppler distribution. The instrumental broadening, however, obeys an entirely different law which causes it to be the predominant effect at great distances from the maximum. The correction (*b*) can therefore be calculated as practically due only to instrumental width. This was done by means of Airy's formula, with the experimentally determined reflectivity of the etalon plates. The result of this calculation was a shift of component (3) by 0.0005 cm.^{-1} towards (2).

For the calculation of effect (*a*), the shape of the line was assumed to be that given by pure Doppler effect, but with the use of the experimentally determined half-value width. The influence of instrumental broadening was thus treated as an empirically determined increase of the Doppler width, corresponding to a higher 'effective' temperature. Its value was found from the measured half-value width

$\Delta\nu = 0.078 \text{ cm.}^{-1}$ of component (1). Allowance for the influence of (4) on the width of (1) gave the effective half-value width of a single line as 0.072 cm.^{-1} . With the use of this value, it was found that (4) shifts the maximum of (1) by 0.0020 cm.^{-1} , and that (3) *B* causes a shift of 0.0014 cm.^{-1} of the peak of (3) *A*. When this correction was calculated, a provisional, uncorrected value of the term shift x had to be used at first, and the resulting, corrected value of x was used in a second calculation (with very nearly the same result), in a process of successive approximation. Also the assumption of the shift of the term $3S_{\frac{1}{2}}$ by $x/3$ enters into this correction. If the latter shift were assumed to be zero, the result would only be changed by 0.0004 cm.^{-1} .

From the wave-number difference $(3) - (1) = 0.2209 \text{ cm.}^{-1}$ in the Dirac term diagram $[2P_{\frac{3}{2}} - 2P_{\frac{1}{2}} - (3P_{\frac{1}{2}} - 3D_{\frac{3}{2}})]$ and the observed value 0.1879 cm.^{-1} of the distance of the peaks (3) and (1), the shift of the $2S_{\frac{1}{2}}$ term is found as

$$\begin{aligned} x = 2P_{\frac{1}{2}} - 2S_{\frac{1}{2}} &= 0.2209 - 0.1879 + 0.0020 + 0.0014 + 0.0005 \\ &= 0.0369 \text{ cm.}^{-1}. \end{aligned}$$

7. THE SPACING OF THE MAIN PEAKS (1) AND (2)

If the assumption of the shift of the $2S_{\frac{1}{2}}$ level is accepted, the component (2) must be regarded as a blend of two lines which are just not resolved and whose intensities are in the ratio 5:2. In a blend of this kind, the position of the maximum will depend quite critically on the individual line shape, and the correction of the type (a) considered in the preceding section will be large. Also in the measurement of a strongly unsymmetrical blend it is doubtful whether the crosswire of the travelling microscope is set on the maximum or the centre of gravity. Since further the position of only one component of the blend, namely, (2) *B*, depends on the position of the level $2S_{\frac{1}{2}}$, it is evident that the measurement of the spacing of the peaks (2) and (1) is not an accurate way of determining the position of this level.

In order to test if the measured spacing $(2) - (1)$ is at least compatible with the value 0.037 cm.^{-1} of the term shift, the peak-to-peak distance was calculated from this value with the use of the methods discussed. The result was 0.323 cm.^{-1} . The measured value 0.321 is very little, but probably significantly, lower. The assumption that the cross-wire was set somewhere between the maximum and the centre of gravity of the blend would account for the difference.

8. THE FAINT COMPONENT (5) AND THE POSITION OF THE TERM $3S_{\frac{1}{2}}$

In very long exposures, a further line appeared next to component (1), towards increasing wave-length. It is due to the transition $2P_{\frac{1}{2}} - 3S_{\frac{1}{2}}$. Indications of it have been observed by Heyden (1937) who was, however, not able to resolve it from the 45 times stronger line (1).

On our plates, this line can be distinctly seen and measured with the travelling microscope, but the resolution is not complete, and the photometer tracings (see figure 3c, plate 12) show it only as a hump in the wing of the overexposed line (1),

not as a maximum of intensity. Our measurements of this component (5) which have been briefly reported before (Kuhn & Series 1948), were carried out in the earlier stages of this work, when the technique of measurement was not as good as in the later experiments. Partly for this reason and partly because of incomplete resolution, the results show a large scatter and are to be considered as provisional. As an average of 17 fringes measured in 8 exposures, the distance of component (5) from (1) was found to be $0.138 \pm 0.006 \text{ cm.}^{-1}$.

The Dirac term diagram (figure 5a) gives the value 0.144 cm.^{-1} for the wave-number difference of these two components. Allowance for the shift of the peak of (1) due to component (4) reduces this value by 0.002 cm.^{-1} , and allowance for instrumental intensity distribution of (1) may reduce it further by an amount y . Any upward displacement z of the term $3S_{\frac{1}{2}}$ is then found as

$$z = 0.144 - 0.138 - 0.002 - y = 0.004 - y \pm 0.006 \text{ cm.}^{-1}.$$

Even if y is assumed to be negligible, our measurements make an upward shift of more than 0.01 cm.^{-1} unlikely. It is hoped that further measurements under improved conditions will give a more definite result.

9. THE TEMPERATURE OF THE EMITTING ATOMS

An estimate of the temperature of the radiating atoms was made from the half-value width of the line (1), measured on photometer tracings. The width was corrected in two respects, first for the influence of the minor component (4), secondly for the instrumental contribution to the width of the line. The remaining width was treated as a pure Doppler effect, and leads to an upper limit for the temperature, since it includes all the minor causes of broadening, such as hyperfine structure, radiation width, Stark effect and imperfections of the etalon.

The first correction is small and was made by numerical calculation. The second correction was made by means of a formula given by Minkowsky & Bruck (1935). The calculation leads to the value 55° K for the temperature under standard discharge conditions. Under the most favourable conditions (lowest current, and the cooling bath at reduced pressure) the value 40° K was obtained. Though no great accuracy is claimed for these figures, they show that other causes of broadening, such as Stark effect, must have been comparatively small.

10. DISCUSSION OF THE RESULTS

Before the results of this work can be compared with other experiments or with theory, it is necessary to discuss various possible systematic errors, as far as they have not been treated in earlier sections.

(a) *Non-uniform illumination.* A decrease of the illumination of the slit image towards the ends could not be avoided. If the slope of the intensity 'envelope' of the fringes had been too great, it could have displaced the maxima. The good agreement of the values of the interval (3) - (1) derived from different orders proved that this effect was negligible. For the interval (2) - (1), the average for the second

order is slightly different from that for the third and fourth orders. This might well be due to the complex structure of the line (2) which could cause measuring errors depending on the dispersion (see § 7).

(b) *Overlapping effects.* The influence of overlapping lines of light hydrogen was specially investigated and excluded by suitable choice of the spacer, as explained in § 4. The possibility that a faint, unknown line of the molecular deuterium spectrum might cause errors by overlapping has been mentioned by other authors. But the intensities of the molecular and atomic spectra of hydrogen depend on the current density in an entirely different way, and since the measured spacings show no dependence on the current density, this source of error can be excluded. Also the variation of deuterium concentration would have revealed any effect of this kind, for similar reasons. The same argument applies to helium bands which have a much higher excitation potential than hydrogen atomic lines and whose intensity depended on the conditions of the discharge in a different way. Another line which might possibly have overlapped was the helium spark line 6560 Å. Subsidiary experiments, however, showed that even much stronger lines of the helium spark spectrum were very weak under the conditions of our discharge.

Very few actual impurities could persist in a discharge at the temperature of liquid hydrogen, in a gas which was circulated through charcoal cooled with liquid air, and they would have revealed their influence by erratic results.

The use of different spacings of the etalon plates, with perfectly concordant results, was a further safeguard against errors caused by the overlapping of unknown lines.

(c) *Self-absorption.* The term $2S_{\frac{1}{2}}$ is known to be metastable, and self-absorption might occur in lines connected with this level, namely, (2) *B* and (3) *A*. This would alter the relative intensities and cause a shift in the blend (2). But any shift in the line (3) would be only of the order of magnitude of the correction applied to this line on account of the influence of (3) *B*. No evidence for self-absorption was, in fact, found. The concentration of metastable atoms must certainly increase with current density, but this was found to affect neither the positions nor, as far as could be judged, the intensity ratio of the lines.

Owing to the small energy difference between the metastable state and the state $2P_{\frac{1}{2}}$, metastable hydrogen atoms are probably extremely sensitive to impacts. This would account for failure of attempts to prove their existence by self-absorption in discharge tubes (Ornstein, Zernicke & Snoek 1928). In our experiments, with low current densities and the presence of helium, self-absorption would not have been expected.

(d) *Stark effect.* Stark effect in discharges, caused either by the local fields of the ions or by the potential drop in the discharge tube, can have noticeable effects on the spectra of light atoms, especially hydrogen. It would cause broadening, shifts and changes of relative intensities of the lines. Theoretical estimates such as those carried out by Heyden (1937) show that these effects are negligible under the conditions of our discharge. Independently, our experimental results provide direct evidence against any appreciable influence of electric fields; the intensity ratio of the components, which should respond strongly to Stark effects, were found to agree

essentially with theory. Change of current density and of deuterium content strongly affect ion density and potential drop, and neither of these changes was found to affect the position of the lines. Also effects of electric fields would be different in the different parts of the discharge tube, but in fact light coming from different parts of the cross section of the discharge tube gave the same positions of the components. In one experiment, the polarities of the electrodes were exchanged, so that the light was taken from the inside of a hollow cathode; but again, the measurements gave the same results.

(e) *Accuracy of the result.* The shifts of the lines (1) and (3) resulting from the influence of the faint, unresolved components (4) and (3)*B* respectively were estimated on the basis of a half-value width of 0.072 cm.^{-1} . If, instead, the half-value width had been assumed as 0.065 or 0.078 cm.^{-1} , the distance (3) – (2) would have been changed by 0.0005 cm.^{-1} in either direction. (These assumptions represent the extreme limits of error in the estimate of the half-value width.) To allow for inaccuracies caused by the simplified method of calculating the corrections, this estimate of error may be doubled. Compared with the resulting value of $\pm 0.001 \text{ cm.}^{-1}$, representing the uncertainties in the corrections, the statistical error of the measurements is small. By adding twice the value of the standard deviation ($2 \times 0.0003 \text{ cm.}^{-1}$), we arrive at the estimated limits of error of ± 0.0016 , and the shift of the term $2S_{\frac{1}{2}}$ is found as

$$0.0369 \pm 0.0016 \text{ cm.}^{-1}.$$

In a preliminary publication (Kuhn & Series 1948), the value of this shift derived from earlier experiments was reported as $0.043 \pm 0.006 \text{ cm.}^{-1}$; this was in disagreement with the preliminary result of Lamb & Retherford (1947) which was given as 0.033 cm.^{-1} . Since then these authors have improved their methods and have given as their final results, for both light and heavy hydrogen, $0.0354 \pm 0.0002 \text{ cm.}^{-1}$. Our new lower value agrees with this new raised value of Lamb & Retherford (1949) within the accuracy of our experiments. The difference between our new and our preliminary value, which is just inside the stated limits of error of the latter, is partly due to imperfections in our early technique which led to a much greater statistical error, and partly to the corrections which had not then been investigated.

(f) *Comparison with theory.* At the time when Lamb & Retherford's first results were published, a paper by Bethe (1947) appeared, in which he calculated that the interaction between the atom and the radiation field should result in a term shift of the order of magnitude of that which had been observed. Since then, a number of authors, treating this interaction in different ways, have arrived at a more accurate theoretical value for the shift of the $2S_{\frac{1}{2}}$ term in hydrogen. Kroll & Lamb (1949) quote the value 1052 Mc./sec. , while French & Weisskopf (1949) give 1051 Mc./sec. , and indicate that Schwinger's new formulation of quantum electrodynamics gives the same results. This theoretical value (0.0351 cm.^{-1}) agrees fairly closely with Lamb & Retherford's later value, and with our spectroscopic value, though it is slightly below both. In this comparison, we have neglected the extremely small shifts of the P and D terms to which the theory leads.

The theories predict a shift of the term $3S_{\frac{1}{2}}$ of $(\frac{2}{3})^3$ times that of the term $2S_{\frac{1}{2}}$, i.e. of 0.0104 cm.^{-1} . Our present, provisional, measurement of the term $3S_{\frac{1}{2}}$ would indicate that the shift is less than this.

The authors wish to thank Professor Lord Cherwell for his interest in the work and Professor F. Simon for the supply of liquid hydrogen. The purchase of the large etalon and other optical equipment was made possible by Professor D. A. Jackson's generous benefactions to the Laboratory.

REFERENCES

- Bethe, H. A. 1947 *Phys. Rev.* **72**, 339.
 Dirac, P. A. M. 1928 *Proc. Roy. Soc. A*, **117**, 610 and **118**, 351.
 Drinkwater, J. W., Richardson, Sir O. & Williams, W. E. 1940 *Proc. Roy. Soc. A*, **174**, 164.
 Du Mond, J. W. M. & Cohen, E. R. 1949 *Rev. Mod. Phys.* **21**, 651.
 French, J. B. & Weisskopf, V. F. 1949 *Phys. Rev.* **75**, 1240.
 Goudsmit, S. & Uhlenbeck, G. E. 1925 *Naturwissenschaften*, **13**, 953.
 Hansen, G. 1925 *Ann. Phys., Lpz.*, **78**, 558.
 Heisenberg, W. & Jordan, P. 1926 *Z. Phys.* **37**, 263.
 Heyden, M. 1937 *Z. Phys.* **106**, 499.
 Kroll, N. M. & Lamb, W. E. 1949 *Phys. Rev.* **75**, 388.
 Kuhn, H. & Series, G. W. 1948 *Nature*, **162**, 373.
 Kupper, A. 1928 *Ann. Phys., Lpz.*, **86**, 511.
 Lamb, W. E. & Retherford, R. C. 1947 *Phys. Rev.* **72**, 241.
 Lamb, W. E. & Retherford, R. C. 1949 *Phys. Rev.* **75**, 1325.
 Mack, J. E. & Barkowsky, E. C. 1942 *Rev. Mod. Phys.* **14**, 82.
 Minkowski, R. & Bruck, H. 1935 *Z. Phys.* **95**, 299.
 Ornstein, L. S., Zernicke, F. & Snoek, J. L. 1928 *Z. Phys.* **47**, 627.
 Pasternack, S. 1938 *Phys. Rev.* **54**, 1113.
 Pauli, W. 1927 *Z. Phys.* **43**, 601.
 Sommerfeld, A. 1916 *Ann. Phys., Lpz.*, **51**, 1.
 Sommerfeld, A. & Unsöld, A. 1926 *Z. Phys.* **36**, 259.
 Williams, R. C. 1938*a* *Phys. Rev.* **54**, 558.
 Williams, R. C. 1938*b* *Phys. Rev.* **54**, 568.
 Williams, W. E. 1942 *Rev. Mod. Phys.* **14**, 94.
 Wood, R. W. 1921 *Phil. Mag.* **42**, 729.

The aerodynamic drag of grassland

BY F. PASQUILL, *Meteorological Office, London*

(Communicated by O. G. Sutton, F.R.S.—Received 8 February 1950)

An improvised drag-plate apparatus, on the principle of that used by Sheppard (1947) on a concrete surface, but suitably modified in design, has been used for exploratory measurements of the aerodynamic drag of grassland. The grass cover was variable (1 to 15 cm. in height), and measurements were made at a number of positions (not simultaneously) in order to obtain an approximate representative value of the drag over a considerable area. Wind velocities were in the region of 500 cm./sec. and, judged in terms of the Richardson number, effectively adiabatic conditions of flow prevailed.

The drag (τ_0) and the simultaneous vertical distribution of wind velocity (u_z) up to a height of 2 m. were found to be expressible in terms of the law well established in the laboratory, i.e.

$$u_z = \frac{1}{k} \sqrt{\left(\frac{\tau_0}{\rho}\right)} \log_e \left(\frac{z-d}{z_0}\right),$$

with k (von Kármán's constant) = 0.37, d (the zero plane displacement) = 8 cm. and z_0 (the roughness parameter) = 0.66 cm. For reasons which are discussed the drag measurements are regarded as approximate, and the close agreement of the numerical value of k with the laboratory value of 0.4 is probably fortuitous. However, the general consistency achieved in this preliminary application suggests that the technique could be profitably developed for a critical investigation of the relation between drag and wind profile and its dependence on atmospheric stability.

In a brief discussion of previous work some evidence is now provided for the validity of the Reynolds formulation of the turbulent shearing stress. Attention is drawn to the application of the present results in treatments of the diffusion of matter in the lower atmosphere.

1. INTRODUCTION

The relation between the drag of a surface and the distribution of fluid velocity in the boundary layer is of fundamental importance in problems of turbulent flow and plays an important part in treatments of turbulent diffusion in the lower atmosphere. Experimental researches by Nikuradse (1933) and Schlichting (1936) on fully developed turbulent flow through pipes and over flat plates with artificially roughened surfaces have established the following law, written here in the form customary in meteorological application:

$$u_z = \frac{1}{k} \sqrt{\left(\frac{\tau_0}{\rho}\right)} \log_e \left(\frac{z-d}{z_0}\right). \quad (1)$$

In this expression,

u_z = fluid velocity at a distance z from the pipe wall or plate, measured from the base of the roughness elements,

τ_0 = shearing stress per unit area of the pipe wall or plate,

ρ = fluid density,

d = a 'zero plane displacement' associated with the finite size of the roughness elements,

z_0 = a constant for each type of roughness, in meteorological application termed the roughness parameter,

k = a dimensionless constant = 0.4.

An equation of this form has been derived by Prandtl and von Kármán on theoretical grounds (see Brunt 1939, pp. 244–247), and the constant k is usually called von Kármán's constant.

The above equation was first applied to the distribution of mean wind velocity near the earth's surface, in adiabatic conditions of flow, by Prandtl (1932), but while the functional form of the equation has since been widely observed to hold in such conditions (see Deacon 1949), the only direct evidence for its explicit validity is that obtained by Sheppard (1947) for a rather unnatural surface. Sheppard measured the drag exerted by the wind on a horizontal plate, floating in oil, under torsional control, and exposed in such a fashion as effectively to form part of an extensive concrete area of smoothness approximating to that of the floating plate. From simultaneous measurements of the vertical profile of wind velocity over the concrete surface it was demonstrated that an equation differing only slightly (see later) from that given in (1) was satisfied, with $k = 0.46$, a value which, in view of the unstable atmospheric conditions then existing, was considered to be in good agreement with that obtained in the laboratory.

The assumption of the explicit validity of the above law for airflow over grassland in adiabatic conditions, i.e. with negligible vertical gradient of potential temperature, with also the assumption that the eddy diffusivities for momentum and matter are identical, has been indirectly justified by recent satisfactory interpretations of data on the dispersal of smoke and vapour from artificial ground level sources (Calder 1949) and of observations on the transport and distribution of natural water vapour (Pasquill 1949*a*). However, direct observations of the relation between drag and wind profile and of the influence thereon of the thermal stratification of the atmosphere have yet to be made over a natural vegetated surface. The present paper describes some preliminary measurements, on grassland, which are first discussed in terms of equation (1) and are then considered briefly in relation to the general problem of aerodynamic drag and turbulent diffusion in the lower atmosphere.

2. INSTRUMENTAL DETAILS AND FIELD TECHNIQUE

The somewhat improvised apparatus used for the measurement of drag was built on the principle of Sheppard's drag-plate, with certain modifications of design to suit the application to a grassland surface. It consists of an aluminium pan floating in water in a slightly larger pan (see figure 1). In field operation the floating pan contains a close-fitting sample of the actual grassland surface, and the outer pan is embedded in the ground at the original position of the sample, so that the floating element of ground is exposed approximately in its natural state. Horizontal movement of the floating pan under the action of wind drag is controlled by a flexible metal strip (actually a pen arm of the type used in a barograph) which is attached to one side of the pan and rigidly held in a small clamping block at its other end. Deflexion of the pan is indicated by a long pointer attached to the opposite side of the pan. The control strip and pointer are shielded from the direct action of the wind by slender cases attached to the outer fixed pan. Adjustment of the height and level of the floating pan is made by altering the quantity and positions of brass weights

carried on a frame attached to the underside of the pan and, if necessary, by altering the quantity of water in the outer pan. This adjustment is made with the control strip unclamped; the latter is subsequently carefully clamped so that there is no resultant vertical force acting on the control strip.

Calibration of the apparatus was carried out in the laboratory by placing small weights in a light paper pan attached to the centre of the floating pan by a hair passing over a light pulley. The relation between the horizontal force applied in a direction normal to the plane of symmetry of the drag-plate assembly, and the deflexion of the pointer was found to be linear within satisfactory limits for a movement of the floating pan over a distance of about $\frac{3}{16}$ in. on either side of the central position, as would be expected in view of the small displacements involved. With the shorter

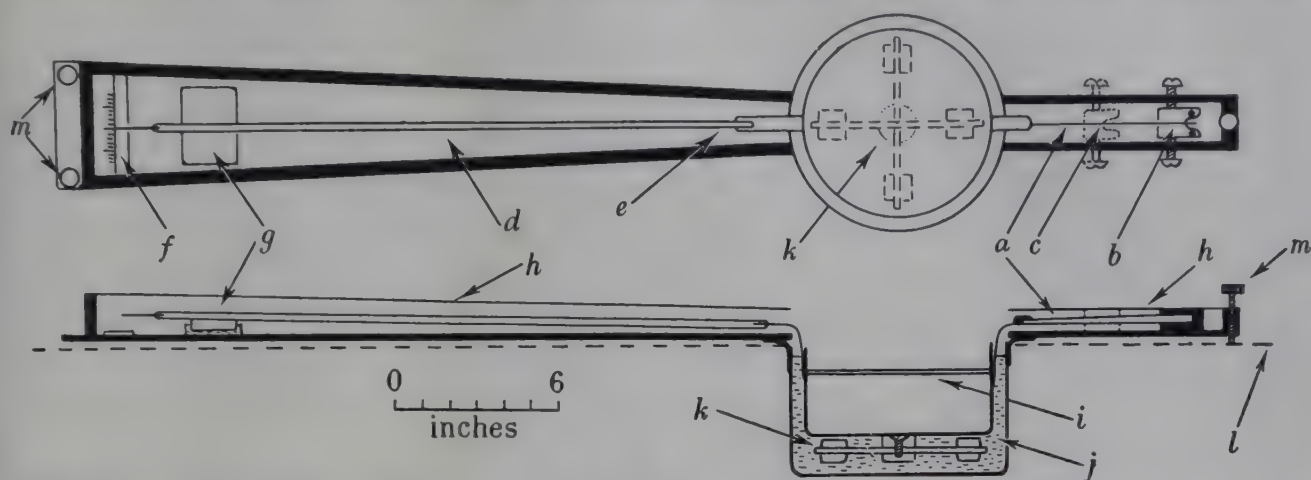


FIGURE 1. Drag plate—plan and section at plane of symmetry. *a*, flexible metal strip rigidly held in clamping block at *b* or *c*; *d*, pointer, with horizontal and vertical adjustment at *e* (not shown); *f*, scale and mirror; *g*, oil bath for damping; *h*, wind-shielding cases with transparent covers; *i*, false bottom to support ground sample (approx. 1 in. thick); *j*, water; *k*, adjustable weights for control of height and level of floating pan; *l*, level of soil surface during field operation; *m*, levelling feet.

control strip (see figure 1), as used in the measurements discussed here, the corresponding overall movement of the pointer was approximately 2 in. Outside this range the linear relation was distorted by surface-tension action in the narrow gap between the two pans. All subsequent measurements were contained within the range of linear response. Calibrations carried out after dismantling and reassembling the apparatus showed an extreme variation of $\pm 3\%$ in the force/deflexion ratio. This degree of reproducibility was probably as good as could be expected from the rather crude mechanical features of the apparatus and was quite acceptable for the preliminary measurements envisaged. Damping of the movement of the plate, effected by a small metal vane carried on the pointer and dipping in an oil bath, was adjusted until critical damping was achieved.

With the shorter control strip a pointer deflexion of 1 in. was produced by a force corresponding to 3.0 dynes/cm.² over the whole horizontal area of the floating pan, acting normal to the plane of symmetry of the apparatus, and full response to an applied force was obtained in approximately 15 sec.

In field operation the drag-plate assembly was set up with its plane of symmetry across wind. The horizontal force acting at the centre of the floating plate, normal to

the axis of symmetry, is then $\tau_0 A \sin \theta$, where A is the area of the floating plate and θ is the angle between the wind direction and the plane of symmetry. A light vane for recording wind direction, a portable distant-reading temperature gradient assembly (described elsewhere, Pasquill 1949*b*) and masts carrying sensitive cup anemometers were set up close to and cross-wind of the drag plate.

Each series of observations was maintained over a period of 10 min. During this period readings of the drag-plate pointer were taken every 5 sec. by an observer lying on the ground on the cross-wind side of the apparatus so as to impose as little as possible disturbance on the air flow. As anticipated the drag-plate deflexions showed considerable fluctuations, though no difficulty was experienced in taking precise readings at the above rate. Before and after each series a draught-proof cover was placed over the drag-plate and zero readings taken. The drag-plate deflexions were mainly in the region of 0.5 in., occasionally 1 in. or more, and were read to 0.01 in. Anemometers were operated simultaneously at heights of 200, 150, 100, 50, 37.5 and 25 cm., and in consecutive runs were interchanged over the levels 200 to 100, 150 to 37½ and 50 to 25 cm. The temperature profile was observed over the same height range, but only the differences over the interval 150 to 37.5 cm. and the absolute value at 200 cm. are employed here. Two 3 min. records of wind direction were taken, centred on 2½ and 7½ min. from the beginning of the observation, and from these a mean wind direction was obtained for each half of the observation period. In reducing the drag-plate results the mean deflexion for each half of the run was associated with the corresponding mean wind direction.

The site of the measurements was a level, well-exposed clayland pasture, with at least 150 yards of similar terrain upwind of the apparatus. The grass cover, however, was variable over distances of the order of yards, and ranged in height from 1 to 15 cm. With a surface non-uniformity on such a scale it was not to be expected that a single siting of the present drag-plate (of effective area approximately 250 sq.cm.) would provide a representative measurement of the drag over a substantial area. The same restriction applies to the vertical wind profile, particularly in the region close to the ground, and, indeed, previous observations on the same site (Pasquill 1950) have demonstrated a local variation of wind distribution. In order therefore to obtain a closer approach to mutually representative drag and wind-profile data three adjacent areas were chosen on which the local grass length was mainly 1 to 5, 5 to 10 and 10 to 15 cm., and measurements were carried out with the drag-plate and wind-profile apparatus set up at each of these three localities in turn. Two 10 min. series of observations were made at each position, the whole process occupying approximately 4 hr.

3. DISCUSSION OF OBSERVATIONS AND RESULTS

The observations are presented in table 1 in the form of mean values for each 10 min. observation, except for the values of wind direction and θ , which correspond to the two 3 min. records of wind direction and are adopted as means for the component 5 min. periods. The value given for τ_0 is the average of the two values based on the mean drag-plate deflexions in these 5 min. periods. Before describing the results ultimately derived it is convenient to summarize at this stage the various limitations

TABLE 1. DRAG AND WIND PROFILE RESULTS

(Girton Pasture, University Farm, Cambridge, 28 July 1949)

local grass length (cm.)	G.M.T.	wind direction (° M)*	θ † (deg.)	u_z/u_{100} at heights in cm. of					measured τ_0 (dynes/cm. ²)	temp. at 200 cm. (° F)	p (mb.)	$\rho \times 10^3$ (g./cm. ³)	$\frac{1}{u_{100}} \sqrt{\frac{\tau_0}{\rho}}$	temp. diff. (150 to 37.5 cm.) (° F)	Ri_{75}
				200	150	50	37.5	25							
1 to 5	1345-55	a 270	82	403	1.110	1.065	0.806	0.776	0.640	0.90	1016	1.20	0.068	-0.73	-0.011
		b 282	70												
	1405-15	a 267	85	478	1.155	1.084	0.824	0.778	0.657	1.44	"	"	0.073	-0.84	-0.008
		b 285	67												
5 to 10	1510-20	a 277	80	483	1.159	1.089	0.834	0.741	0.675	1.53	"	"	0.074	-0.61	-0.004
		b 281	76												
	1530-40	a 238	61	513	1.150	1.084	0.852	0.743	0.678	1.53	"	"	0.070	-0.32	-0.002
		b 262	85												
10 to 15	1650-00	a 289	89	429	1.124	1.072	0.804	0.758	0.625	1.38	"	"	0.079	+0.73	+0.008
		b 290	88												
	1710-20	a 270	72	440	1.141	1.078	0.843	0.757	0.634	1.59	"	"	0.083	+0.72	+0.007
		b 277	79	458	1.140	1.079	0.827	0.759	0.651	1.4	"	"	0.075		-0.002

* a and b refer to the two separate 3 min. records of wind direction.
† Angle between wind direction and plane of symmetry of drag-plate assembly.

imposed by the improvised form of the drag-plate apparatus and the prevailing conditions of ground and weather.

The drag measurements refer to a small area of ground which had been freed from its surroundings. In doing so an artificial gap was produced between the test area and the undisturbed ground. Although nominally this gap was very small, approximately $\frac{1}{2}$ in., in practice it was increased somewhat by the trimming of grass around the edges, a procedure which was necessary to avoid mutual interference of trailing grass blades projecting from the fixed and free edges of the gap. Even so, the artificial irregularity so produced was not greatly dissimilar to irregularities occurring naturally. It will be noted that the cases housing the control strip and pointer projected above the soil surface by about 4 cm. Since, however, these spurious obstacles to flow were on the cross-wind sides of the drag-plate and were for the most part submerged in the natural roughness elements, it seemed reasonable to accept this convenient arrangement for the preliminary measurements rather than to complicate the manipulation and adjustment for wind direction by embedding the cases in the ground. The proximity of the observer provided yet another source of disturbance of the flow conditions over the drag-plate, but as already mentioned this was minimized by the observer lying flat on the ground cross-wind of the apparatus. The limitations imposed by the non-uniformity of surface have been stated in the previous section. Finally, it should be noted that the derivation of the mean drag in the mean direction of the wind directly from the mean drag-plate deflexion requires that the wind direction should be constant, apart from turbulent fluctuations, over the period of measurement. Most of the individual records showed an acceptably constant wind direction over the 3 min. periods, though as will be seen from table 1 the two values obtained in each 10 min. observation differ by as much as 24° in the worst case. For this reason, and as already mentioned, the drag was evaluated separately for the two component 5 min. periods, employing the appropriate values of θ , and averaged to give τ_0 .

Turning now to the data presented in table 1 it is seen that the vertical temperature gradient, as indicated by the temperature differences over the height range 150 to 37.5 cm., varied from a moderate lapse at the beginning to a moderate inversion at the end of the series of measurements. It has previously been demonstrated, however, by Deacon (1949), and by the present writer (Pasquill 1949*a*), that the effects of thermal stratification on processes of turbulent transport near the ground are controlled, not by the vertical temperature gradient alone, but by the ratio of the latter to the square of the vertical wind shear, as in the now familiar Richardson number,

$$Ri = \frac{g}{T} \frac{(\partial T / \partial z + \Gamma)}{(\partial u / \partial z)^2}$$

where T is in $^\circ\text{K}$, Γ is the dry adiabatic lapse-rate, and u is in cm./sec. With high wind velocities and hence, usually, large values of $\partial u / \partial z$, the departure of the above parameter from zero (the value which occurs with zero vertical gradient of potential temperature) may be quite small even with appreciable values of $\partial T / \partial z$. The approximate values of Ri at a height of 75 cm. have been computed from the temperature and wind velocity data appropriate to 37.5 and 150 cm. by assuming a linear variation

of temperature and wind velocity with the logarithm of the height and are reproduced in table 1. Owing to the fairly high wind velocities all values are less than about 0.01 numerically, with an algebraic mean of -0.002 . These values represent departures from 'adiabatic' or 'neutral' conditions of flow which, in terms of the previously mentioned observations of the effect of thermal stratification, are of a very small order. It is thus considered justifiable to bulk the present series of observations as one sample of drag and wind profile data in effectively adiabatic conditions.

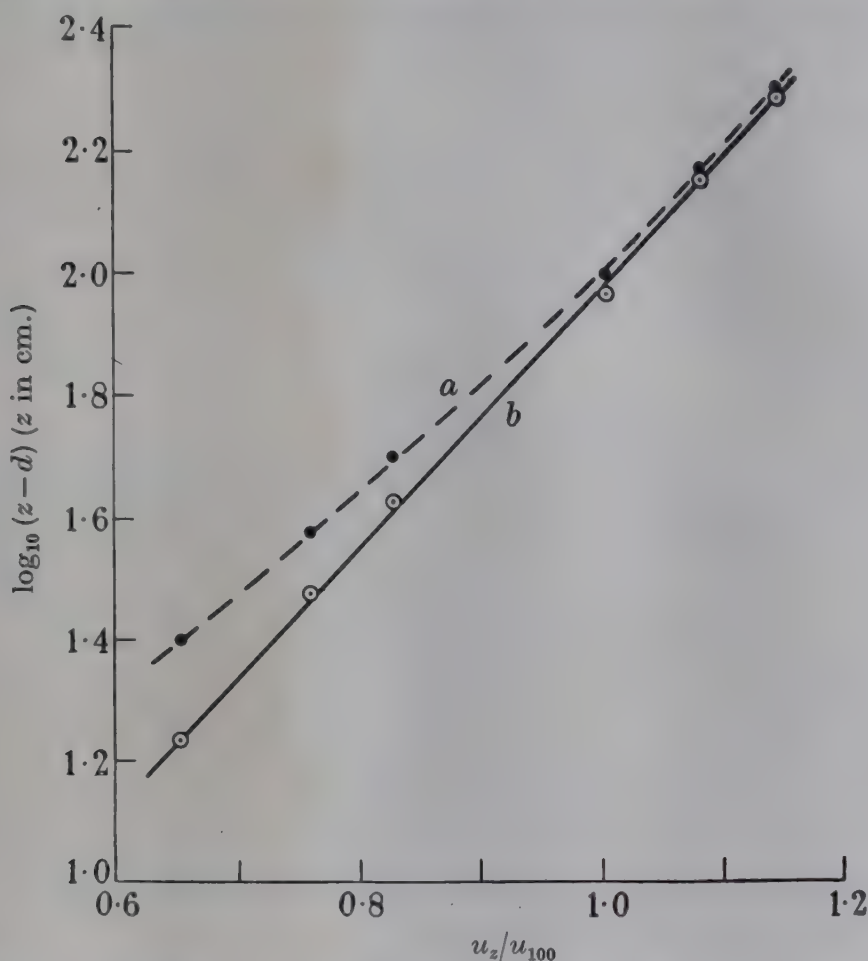


FIGURE 2. Velocity profile over grass 1 to 15 cm. high in effectively adiabatic conditions. Curve *a*, $d = 0$. Curve *b*, $d = 8$ cm.

It will be noted that the individual wind profiles, specified in terms of u_z/u_{100} so as to facilitate comparison of the various observations, exhibit considerable variation. Furthermore, graphs of the individual results, not produced here, display some departures from smooth u , $\log z$ profiles. These features are thought to be mainly a reflexion of genuine local variations in the wind distribution arising from the non-uniformity of the surface. On the other hand, the magnitudes of $\frac{1}{u_{100}} \sqrt{\frac{\tau_0}{\rho}}$, which might have been expected to exhibit wide variations, are surprisingly constant, though there is an indication of a slight systematic increase with increase in local grass length.

Considering now the mean results for the whole series we may note that the wind velocities show a good approach to a smooth u , $\log z$ profile (figure 2*a*), which is, however, slightly concave to the $\log z$ axis. This is attributed to the zero-plane dis-

placement discussed previously (d in equation (1)). Trial plottings show that the best approach to a linear relation between u and $\log(z-d)$ is obtained with $d = 8$ cm., a value which seems reasonably in accordance with the description of the grass cover. The profile is plotted in this form in figure 2*b*, and from this we have

$$\frac{u_z}{u_{100}} = 0.466[\log_{10}(z-8) + 0.182]. \quad (2)$$

From equation (1)
$$\frac{u_z}{u_{100}} = \frac{2.303}{ku_{100}} \sqrt{\left(\frac{\tau_0}{\rho}\right)} [\log_{10}(z-d) - \log_{10}z_0], \quad (3)$$

which is satisfied by (2) if

$$z_0 = 0.66 \text{ cm.} \quad \text{and} \quad \frac{2.303}{ku_{100}} \sqrt{\frac{\tau_0}{\rho}} = 0.466. \quad (4)$$

From table 1

$$\frac{1}{u_{100}} \sqrt{\frac{\tau_0}{\rho}} = 0.075,$$

which with (4) gives

$$k = 0.37,$$

to be compared with the laboratory value of 0.4.

In view of the previously discussed limitations of the present measurements the close agreement of the observed value of k with that established in the laboratory may be to some extent fortuitous. However, the satisfactory nature of this result and the general success achieved in manipulating the preliminary form of drag-plate give reason to believe that the present apparatus and technique could be profitably developed. With improvements in the design of the drag-plate, especially some increase in scale, and further reduction of the various factors possibly contributing to a distortion of the air flow, and with a more rigid restriction to uniformity in ground and atmospheric conditions, it seems likely that the relation between drag and wind profile could be more critically examined and valuable data obtained on the outstanding problem of stability influence.

4. FURTHER DISCUSSION

Although the present measurements are to be regarded as approximate and exploratory in nature they provide data of a character hitherto not available on the subject of the horizontal shearing stress, τ , in the lower atmosphere. This shearing stress, the surface value of which has been measured here, occupies a position of fundamental interest in all considerations of atmospheric turbulence. It is not proposed here to attempt a historical survey of the problem, but it is of interest to consider the present results in the light of certain previous observations and to note the application of the present results in treatments of the diffusion of matter in the lower atmosphere.

Comparisons of special interest are afforded by the work of Sheppard (1947), who measured the drag of an artificial surface effectively forming part of an area of concrete, and whose experimental technique, suitably modified, has been followed here, and by that of Scrase (1930), who evaluated the Reynolds stress ($\tau = -\rho \overline{u'w'}$)

directly from cinematograph records of the movement of light vanes at two heights (1.5 and 19 m.) above downland.

As would be expected from the nature of the surface Sheppard's values of τ_0 are appreciably lower than those given here for grassland. Direct comparison may be made by evaluating a drag coefficient, C_D , defined by the relation

$$\tau_0 = \frac{1}{2}C_D\rho u^2.$$

The magnitude of C_D depends on the level at which u is measured,* and taking this to be 200 cm. we have the following values of C_D :

$$\text{concrete } 2.2 \times 10^{-3}, \quad \text{grassland } 8.6 \times 10^{-3},$$

for wind velocities respectively 421 and 522 cm./sec.

It is also noteworthy that Sheppard related his drag and wind-profile observations in terms of the equation following from the Rossby & Montgomery (1935) form of the mixing length theory of turbulence, with the boundary condition $u = 0$ on $z = 0$, i.e.

$$u = \frac{1}{k} \sqrt{\left(\frac{\tau}{\rho}\right)} \log_e \left(\frac{z+z_0}{z_0}\right), \quad (5)$$

constancy of τ with height being implied. In Sheppard's analysis no difficulty arose regarding the origin of z , because of the smooth nature of the surface, but in the present analysis it is immediately obvious that a reduction of the results in terms of equation (5) would introduce a serious ambiguity in the value of z_0 . This adds to previous indications (see Deacon 1949) that in formulating a general wind-profile relation for rough surfaces it is necessary to introduce two independent arbitrary constants, i.e. z_0 , the roughness parameter, and d , the zero-plane displacement, both of which depend in a complex fashion on the geometrical form of the rough surface.

In Scrase's original paper a value of the shearing stress is given for the height of 19 m., and it is stated that a record taken at 1.5 m. gave a value roughly one-quarter of that at 19 m. Recently, however, Scrase has pointed out in a private communication that the two measurements were of 1 min. duration and separated by an interval of half an hour. Since no simultaneous measurements of wind profile were made the two values of shearing stress cannot strictly be compared, though by appealing to an approximate wind-profile relation and adjusting the observed shearing stresses in accordance with the implied difference in wind velocity on the two occasions, Scrase now finds the 1.5 m. value to be roughly 0.4 time that at 19 m.

The foregoing variation of shearing stress with height, though approximate in view of the indirect nature of the comparison, still constitutes a striking contradiction to the familiar concept of the constancy of τ with height. This constancy has been proved by Ertel (1933) to hold up to a height of about 25 m., and Ertel's proof has been examined and confirmed by Calder (1939). A satisfactory clarification of this

* It follows from equation (1) that the drag coefficient relating to fully developed turbulent flow over a rigid rough surface is otherwise independent of wind velocity and is dependent only on the geometrical form of the surface. The same is not necessarily true of a vegetated surface, since the natural roughness elements may bend over in the wind, so leading to a smoothening of the surface as the wind velocity increases (see Deacon 1949). This effect would result in a diminution of C_D with increase in wind velocity.

discrepancy has yet to be made. One possible explanation, which does not seem to have been considered hitherto, is that the shearing stress observed at 1.5 m., which is presumably dependent primarily on the nature of the surface in the upwind 50 m. or so, is perhaps not strictly comparable with that at 19 m., since the latter would be influenced by surface conditions at a much greater distance. Scrase's observations were made on Salisbury Plain with a level uninterrupted exposure in the upwind kilometre, but even so it is conceivable that the turbulent properties at 19 m. were affected to some extent by the remains of the large-scale eddy structure generated by major irregularities in the surface (houses, trees and topographical features in general) at considerable distances from the site.

A satisfactory understanding of the above points could only be provided by further observation, but for the present purposes it is thought reasonable to regard the value obtained at 1.5 m. as the more reliable figure for the shearing stress in the air immediately over grassland and the more appropriate figure for comparison with the present surface value. Scrase's measurement gives

$$-\rho \overline{u'w'} \text{ at 1.5 m.} = 0.8 \text{ dynes/cm.}^2 \text{ for } u = 468 \text{ cm./sec. at 1.5 m.,}$$

while from the average of the present results

$$\tau_0 = 1.4 \text{ dynes/cm.}^2 \text{ for } u = 494 \text{ cm./sec. at 1.5 m.}$$

Both observations were made over grassland, but it is not possible to say to what extent the surfaces were similar in detail. With this qualification the agreement is nevertheless very reasonable and may be regarded as some evidence for the validity of the Reynolds formulation of the turbulent shearing stress in the atmosphere near the ground.

Finally, it may be noted that the relation of the horizontal turbulent shearing stress to the vertical distribution of wind velocity specifies the eddy diffusivity K of the atmosphere. This provides a basis for treatments of diffusion of matter in the lower atmosphere, a general discussion of which has recently been given by Sutton (1949). In these applications the assumption is made that the eddy diffusivity for matter is identical with that for momentum as derived from the wind profile. Referring to the particular relation expressed in (1) and writing

$$\tau = \tau_0 = \rho K \frac{\partial u}{\partial z},$$

we have

$$K = k \sqrt{\left(\frac{\tau_0}{\rho}\right)} (z-d) = \frac{k^2 u_1 (z-d)}{\log_e \left(\frac{z_1-d}{z_0}\right)}. \quad (6)$$

The eddy diffusivity in adiabatic conditions over a natural rough surface is thus directly proportional to the height above the surface (zero plane displacement being applied) and to wind velocity at some reference height z_1 , and is dependent on the roughness of the surface. For the surface involved in the present measurements of drag the numerical value K at 200 cm. for a wind velocity of 500 cm./sec. at the same height is calculated to be $2.5 \times 10^3 \text{ cm.}^2/\text{sec.}$ The form of K given in (6) has recently been shown by Calder (1949) to lead to solutions of the two-dimensional equation of

diffusion which are in excellent agreement with observations on the distribution of smoke and vapour from artificial sources at ground level, for adiabatic flow over downland. Furthermore, the same form of K leads to a simple expression for the local rate of evaporation from the ground in terms of the vertical distribution of wind velocity and vapour pressure, and this also has been confirmed by observations over grassland in adiabatic conditions (Pasquill 1949*a*, 1950).

Acknowledgement is gratefully made to the Head of the School of Agriculture and the Director of the University Farm, Cambridge, for facilities provided during the course of this investigation, to the Director of the Meteorological Office for permission to publish the results obtained and to members of the Meteorological Office Unit at the School of Agriculture for assistance in the observational work. The writer is much indebted to Dr F. J. Scrase of the Meteorological Office who kindly made available additional and hitherto unpublished information concerning his shearing stress measurements, and to Professor O. G. Sutton, F.R.S., for helpful discussion and criticism during the preparation of this paper.

REFERENCES

- Brunt, D. 1939 *Physical and dynamical meteorology*. Cambridge University Press.
 Calder, K. L. 1939 *Quart. J.R. Met. Soc.* **65**, 537.
 Calder, K. L. 1949 *Quart. J. Mech. Appl. Math.* **2**, 153.
 Deacon, E. L. 1949 *Quart. J.R. Met. Soc.* **75**, 89.
 Ertel, H. 1933 *Met. Z.* **50**, 386.
 Nikuradse, J. 1933 *Forschungsh. Ver. dtsh. Ing.* no. 361.
 Pasquill, F. 1949*a* *Proc. Roy. Soc. A*, **198**, 116.
 Pasquill, F. 1949*b* *Quart. J.R. Met. Soc.* **75**, 239.
 Pasquill, F. 1950 Paper in course of publication in *Quart. J.R. Met. Soc.*
 Prandtl, L. 1932 *Beit. Phys. frei Atmos.* (Bjerknes Festschrift), **19**, 188.
 Rossby, C. G. & Montgomery, R. B. 1935 *Pap. Phys. Ocean. Met., Mass. Inst. Tech.* **3**, no. 3.
 Schlichting, H. 1936 *Ingen. Arch.* **7**, 1.
 Scrase, F. J. 1930 *Met. Off. Geophys. Mem., Lond.*, no. 52.
 Sheppard, P. A. 1947 *Proc. Roy. Soc. A*, **188**, 208.
 Sutton, O. G. 1949 *Atmospheric turbulence*. London: Methuen.

The molecular orbital theory of chemical valency

III. Properties of molecular orbitals

By G. G. HALL AND SIR JOHN LENNARD-JONES, F.R.S.

Department of Theoretical Chemistry, University of Cambridge

(Received 16 December 1949)

The discussion of molecular orbitals and equivalent orbitals, given in previous papers, is carried a stage further. It is shown that certain molecular properties can be evaluated using either equivalent or molecular orbitals. On the other hand, a study of the changes produced by ionization demonstrates that molecular orbitals have a special significance and that certain energy parameters associated with them are closely related to ionization potentials. For the purpose of this discussion a perturbation theory is developed to deal with the changes produced in molecular systems when disturbed from their normal states.

1. INTRODUCTION

In two earlier papers (Lennard-Jones 1949*a, b*), subsequently referred to as parts I and II, equations were obtained for the orbitals to which electrons can be assigned in a molecule in its normal, unexcited state. These equations are as general as can be obtained subject to the assumption that the electronic wave function of a molecule can be expressed as a determinant of mutually orthogonal orbitals, each involving the co-ordinates of one electron only. This form of wave function is the one normally used, because it conforms to the Pauli exclusion principle and at the same time satisfies the principle of the identity of electrons.

It was shown that there is a certain arbitrariness in the choice of the orbitals used in the determinantal wave function, but that two particular types have special significance. These types are most easily defined when a molecule has properties of symmetry. If the orbitals are chosen so as to belong to the various irreducible representations of the symmetry group to which the molecule belongs, they are called *molecular orbitals*. On the other hand, it was shown that another set of orbitals can be chosen equally well. These are called *equivalent orbitals*, for the members of a set have the property that they are identical with each other, except for orientation and position in space. It is the object of this paper to consider which properties of the molecule can be expressed in terms of either type of orbital equally well and to examine the special significance of molecular orbitals.

It appears that molecular orbitals provide a convenient description when the *extensive* properties of a molecule are under consideration; by extensive properties we mean those properties which involve the molecule as a whole. Spectroscopic properties, such as excitation or ionization by light, are of this type. It is shown in this paper that there is a close relation between the ionization potential of a molecule (or, more accurately, the vertical ionization potential) and an energy parameter associated with a molecular orbital. Equivalent orbitals, on the other hand, have advantages in dealing with the *localized* properties of molecules. These include charge distribution in particular bonds or the spatial arrangement of one localized bond

relative to another. The significance of equivalent orbitals will be discussed in part IV. It is satisfactory to find that both the extensive and localized properties of molecules can be dealt with by the same mathematical formulation.

2. TRANSFORMATION PROPERTIES OF ORBITALS

The wave function of the electrons in a molecule is most conveniently expressed as a determinant Φ of mutually orthogonal orbitals, each involving the co-ordinates of one electron only. For this type of wave function equations can be obtained for the orbitals by the use of the variation method. For the special case of states of the molecule in which there are two electrons of opposite spin in each occupied orbital, these equations were shown in part I to be*

$$(H + V'_n - E_{nn}) \psi_n = \sum_{m \neq n} (G_{mn} + E_{mn}) \psi_m, \quad (2.01)$$

in the notation of part II, namely,

$$\begin{aligned} V'_n &= \sum_m 2G_{mm} - G_{nn} \\ &= V - G_{nn}, \end{aligned} \quad (2.02)$$

$$G_{mn} = \int \bar{\psi}_m (1/r_{12}) \psi_n dx_2, \quad (2.03)$$

$$\begin{aligned} E_{mn} &= H_{mn} + \sum_l \{2(lm | G | ln) - (lm | G | nl)\} \\ &= H_{mn} + J_{mn} - K_{mn}^a. \end{aligned} \quad (2.04)$$

The corresponding value of the total electronic energy can then be written as

$$E = \sum_n (E_{nn} + H_{nn}). \quad (2.05)$$

It was also pointed out that the orbitals ψ_n are not uniquely defined in this treatment, but could be subjected to an arbitrary unitary transformation. This transformation gives another set of orbitals χ_n which are related to the orbitals ψ_n by the equations

$$\chi_n = \sum_m T_{nm} \psi_m, \quad (2.06)$$

$$\psi_n = \sum_m T_{nm}^{-1} \chi_m = \sum_m \bar{T}_{mn} \chi_m. \quad (2.07)$$

For any such orbitals a set of equations was derived, formally similar to those for ψ , viz.

$$(H + v'_n - e_{nn}) \chi_n = \sum_{m \neq n} (g_{mn} + e_{mn}) \chi_m, \quad (2.08)$$

$$E = \sum_n (e_{nn} + h_{nn}). \quad (2.09)$$

In these equations the various terms g_{mn} , etc., are functions of the orbitals χ_n in the same way that the terms G_{mn} , etc., are of the orbitals ψ_n . The relation (2.06) can

* In this paper, unless otherwise indicated, all summations are over the occupied orbitals 1, 2, ..., N .

then be used to find the connexion between the various terms of these equations and those of (2.01), and leads to the results

$$g_{mn} = \sum_{ij} T_{im}^{-1} G_{ij} T_{nj} = \sum_{ij} \bar{T}_{mi} T_{nj} G_{ij}, \quad (2.10)$$

$$h_{mn} = \sum_{ij} T_{im}^{-1} H_{ij} T_{nj}, \quad (2.11)$$

$$e_{mn} = \sum_{ij} T_{im}^{-1} E_{ij} T_{nj}. \quad (2.12)$$

By means of these equations any result obtained using one type of orbital can be transformed into terms of any other possible set of orbitals.

3. INVARIANT PROPERTIES OF ORBITALS

Since the value of the determinantal function which has been chosen as an approximate wave function for the molecule is unaltered by any transformation of the type shown in (2.06), we should expect any properties of the molecule as a whole, which can be deduced from the orbitals, to be invariants of the transformation. In this section we shall find these invariants and examine their physical significance.

The invariance of the total electronic wave function Φ implies the invariance of the total electron distribution function, for this is defined as

$$\rho = \bar{\Phi}\Phi. \quad (3.01)$$

The various matrix elements of ρ are also invariant. This was shown by Dirac (1931) to be true for atoms. The result holds also for molecules. These functions are

$$\rho(ij) = 2 \sum_n \bar{\psi}_n(x_i) \psi_n(x_j), \quad (3.02)$$

and (2.07) together with (2.06) proves the result

$$\rho(ij) = 2 \sum_n \sum_m \bar{\psi}_n(x_i) \bar{T}_{mn} \chi_m(x_j) = 2 \sum_m \bar{\chi}_m(x_i) \chi_m(x_j). \quad (3.03)$$

It follows also that quantities derived from these functions will be invariants. Thus, in particular, the joint probability densities of any number of electrons will not be changed. The equations of motion can be written entirely in terms of these functions (Dirac 1930), but it seems more profitable to retain the orbitals ψ_n or χ_n as variables in the equations, since much additional physical information can thereby be obtained.

Since the transformation (2.06) is unitary, we may find additional invariants. This transformation not only produces a rotation of the co-ordinate system in the space spanned by the orbitals but also, as in equations (2.10), (2.11) and (2.12), transforms the matrices which depend on the orbitals. Under this transformation, two invariants of each matrix will be the trace and the determinant. The physical significance of the invariance of the trace is immediate. Thus the trace of $2G_{mn}$ is just the function V defined in (2.02), and this is the Coulomb potential at any point of space due to all the electrons. The trace of $(E_{mn} + H_{mn})$ is the total energy (2.05). Further, by considering the kinetic energy and nuclear potential parts of H separately, we find also the invariance of the total electronic kinetic energy and the total electronic potential energy.

The determinants of G_{mn} , H_{mn} and E_{mn} are also invariant, but these do not seem to have direct physical significance; on the other hand, the determinant of the distribution matrix $\rho(ij)$ is the distribution function itself

$$\rho = |\rho(ij)|, \quad (3.04)$$

and its invariance was found more directly above.

The existence of a symmetry group enables us to find a further invariant. The effect of a group operation P on the orbitals of the occupied set can be represented by a matrix P_{ij} , where

$$P\psi_i = \sum_j P_{ij}\psi_j, \quad (3.05)$$

provided, as will usually be the case for the ground state, the set of occupied orbitals completely spans a number of irreducible representations of the group. The dimensions of the matrix P_{ij} will be equal to the number of occupied orbitals. With another set of orbitals as basis, P will be represented by a matrix p_{ij} :

$$P\chi_i = \sum_j p_{ij}\chi_j. \quad (3.06)$$

The connexion between the matrices is found by using (2.07) to transform (3.05), thus

$$P\psi_i = P \sum_k T_{ik}^{-1}\chi_k = \sum_{jl} P_{ij} T_{jl}^{-1}\chi_l, \quad (3.07)$$

therefore

$$P\chi_k = \sum_{ijl} T_{ki} P_{ij} T_{jl}^{-1}\chi_l, \quad (3.08)$$

and comparing this with (3.06), we find

$$p_{kl} = \sum_{ij} T_{ki} P_{ij} T_{jl}^{-1}. \quad (3.09)$$

As above, this implies that the traces of corresponding matrices are equal and hence that the character of the set of occupied orbitals is invariant.

4. MOLECULAR ORBITALS

In the first section we defined molecular orbitals as those orbitals which belong to the various irreducible representations of the group. The great advantage of using such orbitals is that certain integrals vanish identically. Thus, if ψ_m and ψ_n belong to *different* irreducible representations of the group, the quantities H_{mn} , J_{mn} , K_{mn}^{α} are zero and, therefore, E_{mn} is also zero. This property does not necessarily hold if ψ_m and ψ_n belong to the same irreducible representation. The equation (2.12) shows, however, that this property will be retained if such orbitals are suitably adjusted. For T_{nm} is an arbitrary unitary transformation, and can be chosen so that the matrix E_{mn} is completely reduced to diagonal form. The orbitals determined by this T_{nm} will be called *molecular orbitals*.

The relation between this definition and the previous provisional one can be clarified by introducing another type of orbital. In the theory of vibrations it is convenient to introduce 'symmetry co-ordinates' (Rosenthal & Murphy 1936) which belong to the various irreducible representations of the group, for these simplify the equations considerably. The *normal* co-ordinates, which reduce the

equations of motion still further to diagonal form, are a special type of symmetry co-ordinates. In an analogous way we shall define *symmetry orbitals* as orbitals which belong to the various irreducible representations of the group, but for which E_{mn} is not necessarily reduced to diagonal form. Molecular orbitals are then a special form of symmetry orbitals and correspond to the normal co-ordinates of vibration theory. From these definitions it is obvious that equations (2.06) to (2.12) may be used to connect symmetry orbitals and molecular orbitals, or symmetry orbitals and equivalent orbitals.

As a justification for this definition of molecular orbitals, it is necessary to prove that they are also symmetry orbitals. The vanishing of E_{mn} in the case of molecular orbitals enables the equation (2.01) to assume a rather simpler form:

$$(H + V - E_{nn})\psi_n - \sum_m G_{mn}\psi_m = 0. \quad (4.01)$$

The last term may be written (Dirac 1930) in terms of an operator A defined by the equation

$$A\psi_n(i) = -\sum_m G_{mn}\psi_m(i) = -\int \rho(ji) (1/r_{ji}) \psi_n(j) dx_j, \quad (4.02)$$

and it can be shown that this operator is linear, Hermitean, and has the full symmetry of the group. The equations then take the Hamiltonian form

$$(H + V + A)\psi_n = E_{nn}\psi_n. \quad (4.03)$$

Since this Hamiltonian is the same for each of the occupied orbitals, we may interpret the parameters E_{nn} as the eigenvalues of the Hamiltonian in the eigenstate ψ_n . From this we may deduce that the molecular orbitals ψ_n will automatically be orthogonal and also, since the Hamiltonian has the full symmetry of the group, they will belong to the various irreducible representations of the group. Molecular orbitals are then a particular kind of symmetry orbital and the earlier use of the term is justified.

5. VIRTUAL ORBITALS AND PERTURBATION THEORY

In the previous sections we have considered only those properties of a molecule which can be discussed by use of the wave function of the ground state alone. Many of the most interesting molecular properties cannot, however, be discussed in this way, for they involve perturbations of this state. In this section, then, we shall discuss the application of perturbation theory to the equations for the orbitals, with the object of examining the changes which take place during the process of ionization and the significance of these changes in relation to molecular orbitals.

The unperturbed orbitals in this problem are those of the ground state and consequently satisfy (2.01) which are now written

$$(H + V + A)\psi_n - \sum_{m \neq n} E_{mn}\psi_m = E_{nn}\psi_n. \quad (5.01)$$

We may suppose this equation solved for the N occupied orbitals, and the results substituted in V and A . For these particular functions V and A , which represent

definite functions of the spatial co-ordinates, we can regard equation (5.01) as an ordinary eigenvalue equation and find other solutions. This equation, together with the condition

$$E_{mn} = 0 \quad (m \neq n, n > N), \quad (5.02)$$

determines an infinite set of functions ψ_n , of which the first N are the occupied orbitals. By extending the idea of an orbital, we may call the remaining eigenfunctions of (5.01) *virtual molecular orbitals*. The condition (5.02) ensures that these virtual orbitals are uniquely defined. Since they differ from the eigenfunctions of (4.03) by at most a unitary transformation, the occupied and virtual orbitals together form a complete set of functions, and any perturbed orbital which may have to be considered can be expanded in terms of them. Any complete set of functions could be used for this purpose, but this particular choice has been made because virtual molecular orbitals bear some relation to the physical properties of the molecule.

Virtual orbitals have no direct physical significance, but they may be considered as the possible wave functions of an additional electron moving in the field of the molecule, which contributes nothing to this field although it retains its charge. Such an electron would behave like the 'test charge' of electrostatic theory. This interpretation suggests that most, if not all, of the parameters E_{nn} corresponding to virtual orbitals will lie in the continuum. If it should be desirable to avoid the difficulties of dealing with such a continuum, it would be possible formally to consider the molecule as enclosed in a large sphere, for this would have the effect of replacing the continuum by a large number of discrete values.

The effect of adding further terms $C\psi_n$ to equation (5.01), where C is a small perturbing operator, must now be considered. The equation is

$$(H + V' + A' + C')\psi'_n - \sum_{m \neq n} \mathcal{E}'_{mn} \psi'_m = \mathcal{E}'_{nn} \psi'_n, \quad (5.03)$$

where ψ'_n are the perturbed orbitals and V', A', C' and $\mathcal{E}'_{mn}$ are the same functions of ψ'_n as V, A, C and E_{mn} are of ψ_n , while, by taking the scalar product of (5.03) with ψ'_m ,

$$\mathcal{E}'_{mn} = E'_{mn} + C'_{mn}. \quad (5.04)$$

The solutions of (5.03) can be expanded in terms of the set of unperturbed orbitals

$$\psi'_n = \sum_r a_{nr} \psi_r, \quad (5.05)$$

and this gives the equation

$$\sum_r \{(H + V' + A' + C') a_{nr} - \sum_m \mathcal{E}'_{mn} a_{mr}\} \psi_r = 0. \quad (5.06)$$

The scalar product of this with ψ_p gives

$$\sum_r \{(H_{pr} + V'_{pr} + A'_{pr} + C'_{pr}) a_{nr} - \sum_m \mathcal{E}'_{mn} a_{mr} \delta_{pr}\} = 0, \quad (5.07)$$

which is the matrix form of equation (5.03).

Since the operator C is small, we can take it as of the first order of smallness and expand the solution in increasing order of smallness. The coefficients a_{nr} are given by

$$a_{nr} = \delta_{nr} + a_{nr}^{(1)} + a_{nr}^{(2)} + \dots, \quad (5.08)$$

and the parameters by

$$\mathcal{E}'_{mn} = E_{mn} + \mathcal{E}_{mn}^{(1)} + \mathcal{E}_{mn}^{(2)} + \dots, \quad (5.09)$$

for the zero solution is that of the unperturbed molecule. The series for ψ'_n will also affect the functions V'_{mn} , A'_{mn} and C'_{mn} . This effect can be written as

$$V'_{mn} = V_{mn} + V_{mn}^{(1)} + V_{mn}^{(2)} + \dots, \quad (5.10)$$

$$A'_{mn} = A_{mn} + A_{mn}^{(1)} + A_{mn}^{(2)} + \dots, \quad (5.11)$$

$$C'_{mn} = C_{mn} + C_{mn}^{(2)} + \dots, \quad (5.12)$$

where, for example, we have

$$V_{mn}^{(1)} = 2 \sum_r \sum_l \{ (ml | G | nr) a_{lr}^{(1)} + (mr | G | nl) \bar{a}_{lr}^{(1)} \}, \quad (5.13)$$

$$A_{mn}^{(1)} = - \sum_r \sum_l \{ (ml | G | rn) a_{lr}^{(1)} + (mr | G | ln) \bar{a}_{lr}^{(1)} \}. \quad (5.14)$$

It is to be noted that the summation over l includes only occupied orbitals, while that of r is over *all* orbitals, including virtual as well as occupied ones.

Substituting these expansions into (5.07) and equating terms of the same order to zero, we obtain

$$\sum_r \{ (H_{pr} + V_{pr} + A_{pr}) \delta_{nr} - \sum_m E_{mn} \delta_{mr} \delta_{pr} \} = 0, \quad (5.15)$$

$$\sum_r \{ (H_{pr} + V_{pr} + A_{pr}) a_{nr}^{(1)} - \sum_m E_{mn} a_{mr}^{(1)} \delta_{pr} \} + V_{pn}^{(1)} + A_{pn}^{(1)} + C_{pn} - \sum_m \mathcal{E}_{mn}^{(1)} \delta_{pm} = 0. \quad (5.16)$$

The zero-order equation (5.15) is satisfied identically in virtue of the definition of E_{pn} . The first-order equation (5.16) is complicated by the presence of two sets of unknowns $a_{nr}^{(1)}$ and $\mathcal{E}_{mn}^{(1)}$. It can be simplified by using the definition of E_{pr} to give

$$\sum_r E_{pr} a_{nr}^{(1)} - \sum_m E_{mn} a_{mr}^{(1)} + V_{pn}^{(1)} + A_{pn}^{(1)} + C_{pn} - \sum_m \mathcal{E}_{mn}^{(1)} \delta_{pm} = 0. \quad (5.17)$$

In the particular case of $p = n \leq N$ this becomes

$$\begin{aligned} \mathcal{E}_{nn}^{(1)} &= C_{nn} + V_{nn}^{(1)} + A_{nn}^{(1)} + \sum_r (E_{nr} a_{nr}^{(1)} - E_{rn} a_{rn}^{(1)}) \\ &= C_{nn} + E_{nn}^{(1)}. \end{aligned} \quad (5.18)$$

Now, since ψ_n and ψ'_n are both orthonormal sets, we find that

$$\sum_r \bar{a}_{mr} a_{nr} = \delta_{mn}, \quad (5.19)$$

and, using the expression (5.08), this gives the first-order equation

$$\begin{aligned} \sum_r (\delta_{mr} a_{nr}^{(1)} + \delta_{nr} \bar{a}_{mr}^{(1)}) &= 0, \\ a_{nm}^{(1)} + \bar{a}_{mn}^{(1)} &= 0. \end{aligned} \quad (5.20)$$

Thus we have the implicit relation

$$E_{nn}^{(1)} = V_{nn}^{(1)} + A_{nn}^{(1)} + \sum_r (E_{nr} a_{nr}^{(1)} + E_{rn} \bar{a}_{nr}^{(1)}). \quad (5.21)$$

The solution of these equations in explicit terms is difficult, but many of these difficulties are avoided by restricting the orbitals to be molecular orbitals. In this case equation (5.17) becomes

$$(E_{pp} - E_{nn}) a_{np}^{(1)} + V_{pn}^{(1)} + A_{pn}^{(1)} + C_{pn} - \mathcal{E}_{nn}^{(1)} \delta_{pn} = 0 \quad (n \leq N), \quad (5.22)$$

which is similar to the equation obtained by Peng (1941) using a more specialized form of perturbation appropriate to problems arising in the theory of metals. For $p \neq n$ we have, using the expressions (5.13) and (5.14), along with equation (5.20), a set of equations to determine the $a_{np}^{(1)}$. The parameters $E_{nn}^{(1)}$ can then be found from the remaining case $p = n$ of equation (5.22). Once the solution has been found using molecular orbitals, a transformation of the type used in (2.06) and (2.12) can be applied to give the corresponding results for other orbitals.

6. IONIZATION POTENTIALS AND THE SIGNIFICANCE OF MOLECULAR ORBITALS

In this section we shall consider the effect of moving an electron from a molecule and discuss the particular significance of molecular orbitals. The change of energy that occurs may conveniently be divided into two parts, one depending on the interaction of the electrons and the other on the interaction of electrons and nuclei. Not only is the Hamiltonian diminished by all the terms referring to the ionized electron, but the nuclear co-ordinates also change to new equilibrium values. To isolate the purely electronic part, we define the vertical ionization potential as the difference in total electronic energy when one electron is removed and the nuclear separations are kept fixed. Our object is to find an expression for the vertical ionization potential.

To make the problem definite, we assume that an electron with β spin is missing from the z th orbital. The equations of the ionized molecule, by equation (2.01) of II, are

$$\begin{aligned} (\epsilon_n^\alpha + \epsilon_n^\beta) [H + \mathcal{V}'] \psi'_n - \sum_m (\epsilon_m^\alpha \epsilon_n^\alpha + \epsilon_m^\beta \epsilon_n^\beta) [G'_{mn} + H'_{mn} + \mathcal{J}'_{mn}] \psi'_m \\ + \sum_m (\epsilon_m^\alpha \epsilon_n^\alpha \mathcal{K}'_{mn}{}^\alpha + \epsilon_m^\beta \epsilon_n^\beta \mathcal{K}'_{mn}{}^\beta) \psi'_m = 0, \end{aligned} \quad (6.01)$$

where the dashed quantities refer to the ionized molecule. Since only one electron is unpaired, this can be written

$$\begin{aligned} 2(H + V') \psi'_n - 2G'_{zz} \psi'_n \\ - 2 \sum_m [G'_{mn} + H'_{mn} + J'_{mn} - (zm | G' | zn)] \psi'_m + [G'_{zn} + H'_{zn} + J'_{zn} - (zz | G' | zn)] \psi'_z \\ + \sum_m [K'_{mn}{}^\alpha + K'_{mn}{}^\beta - (zm | G' | nz)] \psi'_m - [K'_{zn}{}^\beta - (zz | G' | nz)] \psi'_z = 0 \quad (n \neq z), \end{aligned} \quad (6.021)$$

$$(H + V') \psi'_z - G'_{zz} \psi'_z - \sum_m [G'_{mz} + H'_{mz} + J'_{mz} - (zm | G' | zz)] \psi'_m + \sum_m K'_{mz}{}^\alpha \psi'_m = 0, \quad (6.022)$$

where J'_{mn} , K'^{α}_{mn} , etc., are the same functions of ψ'_n as J_{mn} , K^{α}_{mn} , etc., are of ψ_n . These equations are further simplified by using the definition of A' and E'_{mn} , and become

$$2(H + V' + A')\psi'_n - 2 \sum_m E'_{mn} \psi'_m - 2G'_{zz} \psi'_n + 2 \sum_m (zm | G' | zn) \psi'_m + G'_{zn} \psi'_z - \sum_m (zm | G' | nz) \psi'_m + E'_{zn} \psi'_z = 0 \quad (n \neq z), \quad (6.03)$$

$$(H + V' + A')\psi'_z - \sum_m E'_{mz} \psi'_m - G'_{zz} \psi'_z + \sum_m (zm | G' | zz) \psi'_m = 0. \quad (6.04)$$

Taking the scalar product of equation (6.03) with ψ'_z , we see that the terms cancel in pairs leaving

$$E'_{zn} = 0 \quad (n \neq z). \quad (6.05)$$

Thus, the most general equations for the orbitals of an ionized molecule imply that the unpaired orbital ψ'_z satisfies this condition.

To bring the equations into the form of equation (5.03), we define the operator C' by the equations

$$2C'\psi'_n = -2G'_{zz}\psi'_n + 2 \sum_m (zm | G' | zn) \psi'_m + G'_{zn} \psi'_z - \sum_m (zm | G' | nz) \psi'_m \quad (n \neq z), \quad (6.061)$$

$$C'\psi'_z = -G'_{zz}\psi'_z + \sum_m (zm | G' | zz) \psi'_m. \quad (6.062)$$

The matrix elements of this operator are defined as

$$C'_{rn} = \int \bar{\psi}'_r C' \psi'_n dx. \quad (6.07)$$

They can easily be evaluated and are

$$C'_{rn} = 0 \quad (r \leq N), \quad (6.08)$$

$$= -(rz | G' | nz) + \frac{1}{2}(rz | G' | zn) \quad (r > N; n \neq z); \quad (6.09)$$

$$C'_{rz} = -(rz | G' | zz) \quad (r > N). \quad (6.10)$$

The total electronic energy of the ionized molecule is

$$E' = 1/(2N-1)! \int \bar{\Phi}' \mathfrak{H}' \Phi' d\tau, \quad (6.11)$$

where $\mathfrak{H}'$ is the Hamiltonian of the ionized molecule and is related to that of the whole molecule $\mathfrak{H}$ by

$$\mathfrak{H}' = \mathfrak{H} - H_z - \sum_{j \neq z} \frac{1}{r_{jz}}. \quad (6.12)$$

The determinantal wave function, using the notation of equation (3.01) of I, is

$$\Phi' = \det \{ \psi'_1(1) \alpha(1), \dots, \psi'_{z-1}(N+z-1) \beta(N+z-1), \psi'_{z+1}(N+z+1) \beta(N+z+1), \dots \}. \quad (6.13)$$

The energy has been evaluated in paper I as equation (3.19) and is, in the notation of paper II

$$E' = \sum_n (\epsilon_n^{\alpha} + \epsilon_n^{\beta}) H'_{nn} + \frac{1}{2} \sum_{lm} \{ (\epsilon_l^{\alpha} + \epsilon_l^{\beta}) (\epsilon_m^{\alpha} + \epsilon_m^{\beta}) (lm | G' | lm) - (\epsilon_l^{\alpha} \epsilon_m^{\alpha} + \epsilon_l^{\beta} \epsilon_m^{\beta}) (lm | G' | ml) \}, \quad (6.14)$$

but only one electron is not paired, so this becomes

$$\begin{aligned} E' &= \sum_n 2H'_{nn} - H'_{zz} + \frac{1}{2} \sum_{lm} \{4(lm | G' | lm) - 2(lm | G' | ml)\} - \sum_l \{2(lz | G' | lz) - (lz | G' | zl)\} \\ &= \sum_n (H'_{nn} + E'_{nn}) - E'_{zz}. \end{aligned} \quad (6.15)$$

Equation (6.05), together with (6.08), implies that E'_{zn} is zero ($n \neq z$). From this, the method of §4 applied to equation (5.03) shows that ψ'_z belongs to an irreducible representation of the group and is a symmetry orbital. Hence, by analogy with the definition for a neutral molecule, ψ'_z may be called a molecular orbital of the ionized molecule. Thus, on ionization, the unpaired electron in the molecule cannot be localized but is in a molecular orbital extending over the whole molecule.

For a sufficiently large molecule, the operator C' is small compared with the remaining terms in the equation and can be treated as a perturbation operator. The zero-order solutions of the equations (6.01) are the orbitals of the neutral molecule

$$\psi'_n = \psi_n, \quad \mathcal{E}'_{zn} = E'_{zn} = E_{zn}. \quad (6.16)$$

On this approximation to the orbitals, the expression for the energy (6.15) becomes

$$E' = E - E_{zz}. \quad (6.17)$$

The next approximation to the energy can be found by including the first order corrections. The first-order correction to E_{nn} is given in equation (5.18). To find the first order correction to H_{nn} the expansion of ψ'_n and the definition of H'_{nn} are used

$$\begin{aligned} H'_{nn} &= \int \bar{\psi}'_n H \psi'_n dx \\ &= H_{nn} + H_{nn}^{(1)} + H_{nn}^{(2)} + \dots, \end{aligned} \quad (6.18)$$

which gives

$$H_{nn}^{(1)} = \sum_r (H_{nr} a_{nr}^{(1)} + H_{rn} \bar{a}_{nr}^{(1)}). \quad (6.19)$$

The first-order expression for the energy is

$$E' = E + \sum_n (H_{nn}^{(1)} + E_{nn}^{(1)}) - E_{zz} - E_{zz}^{(1)}. \quad (6.20)$$

This summation can be simplified by using the expressions obtained in equations (6.19), (5.21), (5.13) and (5.14)

$$\begin{aligned} \sum_n (H_{nn}^{(1)} + E_{nn}^{(1)}) &= \sum_n \sum_r \{H_{nr} a_{nr}^{(1)} + H_{rn} \bar{a}_{nr}^{(1)} + 2 \sum_l [(nl | G | nr) a_{lr}^{(1)} + (nr | G | nl) \bar{a}_{lr}^{(1)}] \\ &\quad - \sum_l [(nl | G | rn) a_{lr}^{(1)} + (nr | G | ln) \bar{a}_{lr}^{(1)}] + E_{nr} a_{nr}^{(1)} + E_{rn} \bar{a}_{nr}^{(1)}\}. \end{aligned} \quad (6.21)$$

By rearranging terms and using (5.20), the summation becomes

$$\begin{aligned} \sum_n (H_{nn}^{(1)} + E_{nn}^{(1)}) &= 2 \sum_n \sum_r (E_{nr} a_{nr}^{(1)} + E_{rn} \bar{a}_{nr}^{(1)}) \\ &= 2 \sum_n \sum_r (E_{nr} a_{nr}^{(1)} - E_{rn} a_{rn}^{(1)}) \\ &= 0 \dots \end{aligned} \quad (6.22)$$

Thus, on ionization from the molecular orbital ψ_z , the energy becomes

$$E' = E - E_{zz} - E_{zz}^{(1)}. \quad (6.23)$$

This proves that, for molecular orbitals, the expression for the total energy given by equation (6.17) is correct to within a term small compared with the parameter E_{zz} .

We may summarize this in the form of a theorem, stating that the vertical ionization potentials of a molecule are, to a good approximation, the negative of the parameters E_{nn} in a molecular orbital representation. This theorem is an extension of that proved by Koopmans (1933).

This discussion gives us a new insight into the physical significance of molecular orbitals. For, in § 3, we saw that certain properties of a molecule could be discussed equally well using molecular or equivalent orbitals but, in this section, we have an example of a property which bears a special relation to molecular orbitals alone. Equation (6.16), together with (6.05), shows that, not only is ψ'_z a molecular orbital but that the orbital of the neutral molecule closest to it, ψ_z , is also a molecular orbital. This, then, is our justification for thinking of ionization as the removal of an electron from a molecular orbital. The theorem we have proved emphasizes this by showing that we may regard the parameter E_{nn} as giving the contribution of the second electron, in the molecular orbital ψ_n , to the total energy.

We wish to acknowledge our thanks for the award to one of us (G. G. H.) of a maintenance grant by the Ministry of Education, Northern Ireland, and a studentship by Queen's University, Belfast. We are also indebted to Mr L. A. G. Dresel and Mr J. A. Pople for reading the paper and for their helpful suggestions in the formulation of the last section.

REFERENCES

- Dirac, P. A. M. 1930 *Proc. Camb. Phil. Soc.* **26**, 376.
 Dirac, P. A. M. 1931 *Proc. Camb. Phil. Soc.* **27**, 240.
 Koopmans, T. 1933 *Physica*, **1**, 104.
 Lennard-Jones, Sir J. 1949a *Proc. Roy. Soc. A*, **198**, 1.
 Lennard-Jones, Sir J. 1949b *Proc. Roy. Soc. A*, **198**, 14.
 Peng, H. W. 1941 *Proc. Roy. Soc. A*, **178**, 499.
 Rosenthal, J. E. & Murphy, G. M. 1936 *Rev. Mod. Phys.* **8**, 317.

The molecular orbital theory of chemical valency

IV. The significance of equivalent orbitals

BY SIR JOHN LENNARD-JONES, F.R.S., AND J. A. POPLE
Department of Theoretical Chemistry, University of Cambridge

(Received 16 December 1949)

The theory of the transformation from molecular orbitals to sets of equivalent orbitals is discussed for the general case when there is more than one occupied molecular orbital of given symmetry and more than one equivalent set. The general transformation is worked out for molecules whose component atoms possess inner shells and lone pairs of electrons. The theory is illustrated by reference to some simple molecules such as water and ammonia. Finally, it is shown how the expression for the total energy of a molecule can be divided up in such a way that the interactions between its localized parts are dealt with separately. The significance of lone pairs of electrons in determining the shape of molecules is pointed out.

1. INTRODUCTION

There are alternative ways of describing the electronic structure of molecules in terms of orbitals. These different modes of description lead to the same results when properties such as the kinetic energy or the potential energy of the electrons are calculated. None the less it is necessary to examine these alternative electron configurations and to find their relation to one another. Part III (Hall & Lennard-Jones 1950) has shown what is meant by a molecular orbital and what its value is in discussing the energy levels of a molecule. It was found that this mode of description is useful in dealing with spectroscopic properties.

This paper deals with equivalent orbitals and shows that they permit a closer understanding of the relative positions of electrons in molecules. We have no means of observing the distribution of one electron relative to another. All that experiment can determine is the average distribution of the electrons in a molecule in ordinary three-dimensional space, and this will conform to the symmetry of the nuclear framework. The probability that one electron, no matter which, will be found in a given element of space, a second one in another element of space, can be expressed only as a function of six co-ordinates. Though we cannot observe this quantity, we still require to know what it is, if we are to calculate the electrostatic interaction of electrons. This quantity will affect the energy of the molecule as a whole and so be involved indirectly in a property which is observable. So it can only be inferred *a posteriori* whether an assumed probability function is correct.

Though molecular orbitals and equivalent orbitals are equally valid ways of representing electrons in molecules and both lead to the same answer in calculating electrostatic energy of the electrons, yet one method of calculation is more easily interpreted than the other. The expression for the interaction of the electrons appears in *two* parts, both of which are invariant under any orthonorm transformation of the orbitals. But they each contain a common part which can be cancelled. This common part is not an invariant; and so the parts which are left, when the common

part has been taken away, are not invariant. One of these gives the electronic interaction of electrons in charge distributions corresponding to the orbitals, as in a straightforward problem of electrostatics. It implies repulsion and is usually called the Coulomb contribution. The other part, often called the 'exchange' contribution, has no counterpart in classical theory and indeed occurs with the opposite sign to the Coulomb part. The respective values of the Coulomb and exchange contributions depend on the form chosen for the orbitals. It turns out that the exchange part is small when equivalent orbitals are used, and the Coulomb part thus gives the major contribution to the energy. The interaction of electrons in molecules may thus be interpreted as due largely to the repulsion of charges distributed in equivalent orbitals. This version of electron interaction is illuminating, not only as indicating the meaning of localized distributions, but also as an aid to understanding the effect of one such distribution (or bond) on another.

This point of view is useful, too, in considering the role played by lone pairs of electrons in molecules. These also may be regarded as in a sense directed. An alteration in the angles between bonds may change the disposition of lone pairs relative to one another and so change their repulsion. Lone pairs may thus exercise a stabilizing influence on bonds and help to determine the angles between them; lone pairs may, in fact, be more important in determining molecular structure than has yet been recognized.

2. SETS OF EQUIVALENT ORBITALS

One of the most convenient ways of representing the wave function of a system containing many electrons is by a series of products of functions of the co-ordinates of individual electrons. To conform to the Pauli principle, such a series of functions must be antisymmetric in the electron co-ordinates (including spin). The simplest way of ensuring this is by arranging the series in the form of a determinant. The functions of the individual electrons may, for short, be described as *orbitals*. A more general form of wave function would consist of a series of such determinants, and this mode of representation would be as general as could be devised subject to the assumption that the relative co-ordinates of the electrons (r_{ij} , etc.) do not appear explicitly in the wave function. However, in this paper, as in the previous parts, we shall confine ourselves to a single determinantal wave function.

In previous papers (parts I and II, Lennard-Jones 1949*a, b*; part III, Hall & Lennard-Jones 1950) general equations have been derived for a set of orthogonal one-electron functions used to describe the orbitals of a molecule. For the normal state of a molecule containing two electrons of opposite spin in each orbital, the equation for the orbital is

$$\{H + V'_n(x) - E_{nn}\} \psi_n = \sum'_m \{G_{mn}(x) + E_{mn}\} \psi_m. \quad (2.01)$$

These functions ψ_n have been described in part III as *molecular orbitals* when each is a member of an irreducible representation of the symmetry group, to which the molecule belongs, and when

$$E_{mn} = 0 \quad (m \neq n). \quad (2.02)$$

The latter condition fixes orbitals which belong to the same symmetry type. The significance of this condition has been discussed in part III.

By applying an orthonormal transformation to the set of occupied molecular orbitals ψ_n , an alternative set can be obtained, and this may be more convenient in discussing certain properties of the molecule. One such transformation is to a set of equivalent orbitals. Such a set may be defined as consisting of normalized orthogonal one-electron functions, any one of which may be transformed to any other by an operation of a certain subgroup of the symmetry group of a molecule, the set as a whole being invariant under this subgroup. The subgroup will usually be the full symmetry group itself, but in certain cases it may be found necessary to use the extended form of the definition. Methods of obtaining these as linear combinations of molecular orbitals have been given in part II.

Some molecules may be regarded as containing more than one set of equivalent orbitals. Thus in CH_2Cl_2 , for example, there will be two sets of equivalent orbitals, each set having two members. To cover all cases we note that, according to the definition given above, a single orbital may in certain molecules be regarded as an equivalent set with only one member. The molecular orbitals of any molecule, whether possessing symmetry or not, may then be transformed to a set of equivalent orbitals. These may be written as

$$\chi_{a_1}, \chi_{a_2}, \dots, \chi_{a_n}; \quad \chi_{b_1}, \dots, \chi_{b_n}; \quad \chi_{c_1}, \dots, \quad (2.03)$$

The transformation required can be written as in II, equation (3.04),

$$(\chi) = (t)(\psi), \quad (2.04)$$

where (χ) is a column vector containing all the sets of equivalent orbitals.

The transformation of the energy matrix E_{mn} is given by

$$e_{mn} = \sum_{\mu} \bar{t}_{m\mu} t_{n\mu} E_{\mu\mu}. \quad (2.05)$$

In general, the set of molecular orbitals (ψ) may contain members belonging to the same irreducible representation, that is, the same symmetry type, but associated with different energies. If a molecular orbital of this representation is required for the formation of a set of equivalent orbitals, the question arises as to which member is to be used, or which linear combination of those available is to be chosen.

It is difficult to lay down general rules, but examination of a particular case usually suggests the most suitable choice. One way of simplifying the theory is to use separate molecular orbitals for different sets of equivalent orbitals and not linear combinations, for an examination of (2.05) shows that the e_{mn} cross-terms between orbitals of different sets then vanish; that is, the energy matrix e_{mn} consists of square blocks of non-zero terms distributed end to end along the diagonal, each block corresponding to a different set of equivalent orbitals. This may be regarded as the next best arrangement to the purely diagonal form of E_{mn} , appropriate to molecular orbitals. Thus, in molecules possessing no symmetry, it leaves the molecular orbital unchanged and produces no simplification. An alternative criterion that may then be used is to choose linear combinations of molecular orbitals to produce orbitals of maximum localization.

Equivalent sets of orbitals may represent not only localized bonds, but also lone pairs of electrons and atomic inner shells. As localization is only a relative term, it

is not possible to draw a sharp distinction between these three classes, but inner shells and lone pairs are characterized by being concentrated mainly in the neighbourhood of one nucleus, so that they may in practice be represented approximately by ordinary atomic orbitals, or by directed atomic orbitals.

3. THE ENERGY OF A MOLECULE

The total energy of the normal state of a molecule, including the nuclear interaction, has been given in the earlier paper as

$$E = 2 \sum_n H_{nn} + \sum_m \sum_n \{2(mn | G | mn) - (mn | G | nm)\} + \sum_{\alpha \neq \beta} Z_\alpha Z_\beta / r_{\alpha\beta}. \quad (3.01)$$

The summation over italic letters covers the occupied space orbitals, while Greek letters are used for the nuclei. In this formula

$$(lm | G | np) = \iint \bar{\psi}_l(i) \bar{\psi}_m(j) \psi_n(i) \psi_p(j) (1/r_{ij}) dr_i dr_j \quad (3.02)$$

and
$$H_{nn} = \int \bar{\psi}_n(j) H_j \psi_n(j) dr_j, \quad (3.03)$$

where H_j is the operator representing the kinetic energy and the field of the bare nuclei,

$$H_j = -\frac{1}{2} \nabla_j^2 - \sum_\alpha Z_\alpha / r_{\alpha j}. \quad (3.04)$$

In these formulae the functions ψ_n represent molecular orbitals.

It is to be noted that the formula for the energy is exact so long as the molecular wave function is of the determinantal form, described in §2 above. Nothing has been neglected.

We thus have the following contributions to the energy:

(i) $\mathfrak{T} = 2 \sum_n H_{nn}$, which represents the kinetic energy of all the pairs of electrons in the orbitals ψ_n and their potential energy in the field of the bare nuclei;

(ii) $\mathfrak{S} = \sum_m \sum_n 2(mn | G | mn)$, which represents the electrostatic repulsion of the pairs of electrons in the various charge distributions ψ_m, ψ_n ;

(iii) $-\mathfrak{R} = -\sum_m \sum_n (mn | G | nm)$, which represents (apart from the terms $m = n$) the non-diagonal matrix part of the electrostatic interactions of the electrons with the same spin in different orbitals ψ_m and ψ_n ;

(iv) $\mathfrak{N} = \sum_{\alpha \neq \beta} Z_\alpha Z_\beta / r_{\alpha\beta}$, which represents the repulsion of the nuclei.

Now, none of the above contributions is changed when the orbitals are changed from molecular orbitals to equivalent orbitals or, in fact, to any set of orbitals which arise from an orthonormal transformation of the type given in equation (2.04). It would at first sight appear as though no advantage could be gained by a transformation. While the above energy contributions remain the same, we want to find what *interpretation* may be given to their various parts.

We have

$$\mathfrak{J} = 2 \sum_m (mm | G | mm) + 2 \sum'_{m,n} (mn | G | mn), \quad (3.05)$$

where the second summation indicates that m is different from n . In this expression the first term gives twice the energy of pairs of electrons in the *same* orbital ψ_m , and the second gives the energy of pairs of electrons in different orbitals, all pairs of orbitals being counted once only in the summation.

Similarly,

$$\mathfrak{K} = \sum_m (mm | G | mm) + \sum'_{m,n} (mn | G | nm), \quad (3.06)$$

We thus see that the first term of $\mathfrak{K}$ cancels part of $\mathfrak{J}$ and we have

$$\mathfrak{J} - \mathfrak{K} = \sum_m (mm | G | mm) + 2 \sum'_{m,n} (mn | G | mn) - \sum'_{m,n} (mn | G | nm). \quad (3.07)$$

Though $\mathfrak{J}$ and $\mathfrak{K}$ are each invariant under a transformation of the orbitals, the various terms appearing in (3.07) are not. Thus for a transformation to equivalent orbitals, we have

$$\mathfrak{J} - \mathfrak{K} = \sum_{\mu} (\mu\mu | g | \mu\mu) + 2 \sum'_{\mu,\nu} (\mu\nu | g | \mu\nu) - \sum'_{\mu,\nu} (\mu\nu | g | \nu\mu), \quad (3.08)$$

where

$$(\kappa\lambda | g | \mu\nu) = \iint \bar{\chi}_{\kappa}(i) \bar{\chi}_{\lambda}(j) \chi_{\mu}(i) \chi_{\nu}(j) (1/r_{ij}) dr_i dr_j, \quad (3.09)$$

and the χ 's are related to the ψ 's by (2.04).

If we divide the right-hand side of (3.07) into two parts, viz.

$$J = \sum_m (mm | G | mm) + 2 \sum'_{m,n} (mn | G | mn) \quad (3.10)$$

and

$$-K = - \sum'_{m,n} (mn | G | nm), \quad (3.11)$$

we see that J and $(-K)$ correspond to the Coulomb and exchange parts of the energy for molecular orbitals, as described in the introduction. A similar division for equivalent orbitals leads to

$$j = \sum_{\mu} (\mu\mu | g | \mu\mu) + 2 \sum'_{\mu,\nu} (\mu\nu | g | \mu\nu), \quad (3.12)$$

$$-k = - \sum'_{\mu,\nu} (\mu\nu | g | \nu\mu). \quad (3.13)$$

But it is to be noted that K and J are not invariant under an orthonorm transformation of the orbitals, so that we shall not as a rule have $J = j$, or $K = k$, but $J - K = j - k$.

4. EXAMPLES OF EQUIVALENT ORBITALS

The relative importance of the various terms, which appear in the expression for the electrostatic energy of electrons in the above discussion, becomes clearer when particular examples are worked out.

Suppose that a system consists of two particles each moving in one dimension in a field which is symmetrical about some point. Then the wave functions for each particle, that is, the orbitals, are either symmetrical or antisymmetrical about the origin. We will consider the case when the particles occupy different orbitals, one of which is symmetrical and the other antisymmetrical. The symmetrical one will

be labelled ψ_s and the other ψ_a . For non-degenerate levels these functions must be real. For purposes of illustration we will assume that there is some property analogous to spin, so that when this is the same for the two particles, the wave function is antisymmetric in the co-ordinates of the particles. The two-particle probability function will then be $Af(1, 2)$, where

$$f(1, 2) = \begin{vmatrix} \rho(1, 1) & \rho(1, 2) \\ \rho(2, 1) & \rho(2, 2) \end{vmatrix}; \quad (4.01)$$

A is a normalizing factor and the elements of the determinant are given by

$$\begin{aligned} \rho(1, 1) &= \psi_s(1) \psi_s(1) + \psi_a(1) \psi_a(1), \\ \rho(1, 2) &= \psi_s(1) \psi_s(2) + \psi_a(1) \psi_a(2). \end{aligned} \quad (4.02)$$

The function $f(1, 2)$ must obviously be represented in two-dimensional space.

The functions $\rho(1, 1)$, $\rho(1, 2)$, etc., are all invariant under an orthonormal transformation of the functions ψ_s and ψ_a , but when the common parts of $\rho(1, 1)$ $\rho(2, 2)$ and $\rho(1, 2)$ $\rho(2, 1)$ are cancelled, we have

$$\begin{aligned} f(1, 2) &= \{\psi_s^2(1) \psi_a^2(2) + \psi_s^2(2) \psi_a^2(1)\} - 2\psi_s(1) \psi_s(2) \psi_a(1) \psi_a(2) \\ &= F_c - F_e. \end{aligned} \quad (4.03)$$

F_c and F_e are not now invariants separately, although their difference is. F_c may be written

$$F_c = \rho_s(1) \rho_a(2) + \rho_s(2) \rho_a(1), \quad (4.04)$$

and corresponds to a classical interpretation of one particle in each of the continuous distributions ρ_s and ρ_a , allowance being made for either particle to be in ρ_s .

In this example ψ_s and ψ_a correspond to molecular orbitals, for they have symmetry properties characteristic of the symmetry field in which the particles move. Equivalent orbitals are obtained by the transformation

$$\begin{aligned} \chi_1 &= \frac{1}{\sqrt{2}} \{\psi_s + \psi_a\}, \\ \chi_2 &= \frac{1}{\sqrt{2}} \{\psi_s - \psi_a\}, \end{aligned} \quad (4.05)$$

and the expression analogous to (4.03) is

$$\begin{aligned} f(1, 2) &= \{\chi_1^2(1) \chi_2^2(2) + \chi_2^2(1) \chi_1^2(2)\} - 2\chi_1(1) \chi_2(1) \chi_1(2) \chi_2(2) \\ &= f_c - f_e. \end{aligned} \quad (4.06)$$

In this form f_c corresponds to a classical interpretation in terms of continuous distributions ρ_1 and ρ_2 .

The separate parts of $f(1, 2)$ calculated by the two methods are not equal. To illustrate the difference between them we may use the functions appropriate to the two lowest levels of a linear oscillator, for which

$$\begin{aligned} \psi_s &= (\alpha/\pi)^{\frac{1}{2}} e^{-\frac{1}{2}\alpha x^2} \\ \psi_a &= (2\alpha)^{\frac{1}{2}} (\alpha/\pi)^{\frac{1}{2}} x e^{-\frac{1}{2}\alpha x^2}. \end{aligned} \quad (4.07)$$

The function $f(1, 2)$ is plotted in figure 1. This shows maxima at the points $(1/\sqrt{(2\alpha)}, -1/\sqrt{(2\alpha)})$ and $(-1/\sqrt{(2\alpha)}, 1/\sqrt{(2\alpha)})$ and a node along the line $x_1 = x_2$, and demonstrates that the particles tend to be on opposite sides of the origin. This state of affairs arises not because the particles repel each other, for we have not as yet specified their form of interaction, but because of the orthogonal property of the orbitals to which they have been assigned. It may be regarded as a manifestation of the Pauli exclusion principle.

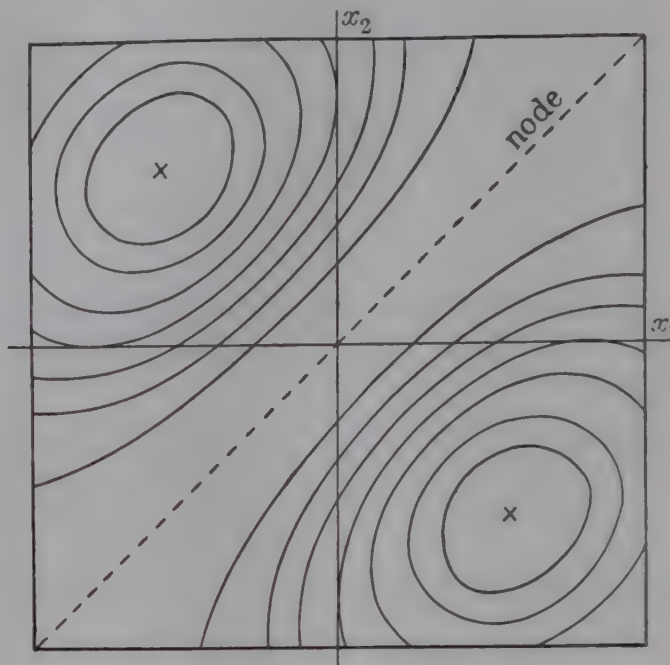


FIGURE 1. The probability distribution function $f(1, 2)$ for a two-particle system in one dimension.

On the other hand, the 'classical' part of $f(1, 2)$ obtained by using molecular orbitals is given by F_c , which is

$$F_c = (2\alpha^2/\pi) (x_1^2 + x_2^2) \exp[-\alpha(x_1^2 + x_2^2)]. \quad (4.08)$$

This is a function only of $(x_1^2 + x_2^2)$ and has no nodes. It gives the same value for the probability of the two particles being on the same side of the origin as on opposite sides. The contours in the (x_1, x_2) plane are concentric circles with centre at the origin. The imperfections of F_c are corrected by f_c , but it is clear that the 'classical' part above does not in itself give a good representation of the probability function.

When equivalent orbitals are used, the 'classical' part of $f(1, 2)$ is

$$f_c = (\alpha/2\pi) \exp[-\alpha(x_1^2 + x_2^2)] \{1 + 2\alpha(x_1^2 + x_2^2) + 4\alpha^2 x_1^2 x_2^2 - 8\alpha x_1 x_2\}. \quad (4.09)$$

This function is shown in figure 2. It produces many of the features of $f(1, 2)$, for it has maxima in the same places and vanishes at two points

$$(1/\sqrt{(2\alpha)}, 1/\sqrt{(2\alpha)}), \quad (-1/\sqrt{(2\alpha)}, -1/\sqrt{(2\alpha)})$$

on the line $x_1 = x_2$. It shows that a classical interpretation of the relative distribution of the two particles in terms of equivalent orbitals gives a good approximation to the actual distribution (figure 1).

Another instructive example is a system of two electrons moving in three dimensions in the field of a central nucleus. If one of these is in an s orbital and the other

in a p orbital we have a situation similar to that in the one-dimensional example given above. The two-particle probability function can no longer be represented in two-dimensional space, but there is the compensating advantage that by taking approximate analytical forms for the orbitals we can calculate the values of the Coulomb and exchange energies for both molecular and equivalent orbitals. The values, using $2s$ and $2p$ functions of the type introduced by Slater (1930) with screening constant Z , are given in table 1.

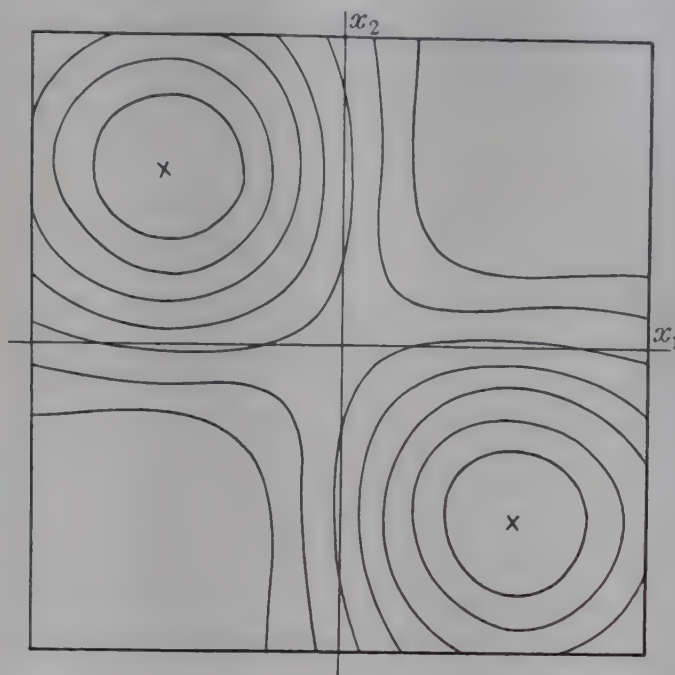


FIGURE 2. The 'classical' part of the probability distribution function (f_c) for the same system as figure 1, using equivalent orbitals.

TABLE 1. ELECTRON REPULSION ENERGY TERMS FOR A $(2s)(2p)$ CONFIGURATION

	Coulomb	exchange	total
molecular orbitals	$0.1816Z$	$-0.0401Z$	$0.1415Z$
equivalent orbitals	$0.1450Z$	$-0.0035Z$	

This example illustrates the great advantage of using equivalent orbitals, for it is evident that a large proportion of the energy arises from configurations in which particles are distributed in different equivalent orbitals. Whereas for molecular orbitals the 'exchange' energy is 28 % of the whole, for equivalent orbitals it is only 2.4 %.

The above analysis applies to the $(2s)(2p)^3P$ state of atoms such as beryllium. There is also a singlet state $(2s)(2p)^1P$ in which the wave function is antisymmetrical in the spins. In this state the distribution function leads to a tendency for both electrons to be on the same side of the nucleus. However, the 3P form has weight $\frac{3}{4}$ in the free-spin state (which corresponds to a state of divalency), so the general conclusions apply to an atom in this configuration and indicate the tendency of such an atom to form bonds in opposite directions.*

* (Added in proof.) Further calculations show that when the number of equivalent orbitals is increased, as in trigonal or tetrahedral distributions, the transformation from molecular (or atomic) orbitals to equivalent orbitals leads to a reduction in exchange energy but the advantage is not so pronounced as in the digonal example.

5. THE INNER SHELLS OF MOLECULES

Having seen the advantage of using equivalent orbitals when the energy of a molecular system is being worked out, we now propose to consider the application of the method to molecules. We shall examine first methods of dealing with the inner electrons and then proceed to a discussion of electrons in lone pairs. Our object is to discover how best to set out the expression for the electrostatic energy of a molecule so as to indicate clearly which are the important contributions. To this end we shall find it useful to ascribe electrons in lone pairs to equivalent orbitals in a way analogous to that used for bonds. It will thus be evident that the representation of lone pairs is dependent on the representation of bonds and an alteration of one implies an alteration of the other.

(a) The inner shells of central hydrides

The method of dealing with inner shells can best be illustrated by discussing some simple examples. In the iso-electronic series of central hydrides CH_4 , NH_3 , H_2O , HF , the K electrons of the central atom are concentrated near one nucleus and are characterized by a larger (negative) value of E_{nn} than any of the other occupied orbitals. They do not, therefore, combine well with other orbitals of different symmetry with smaller values of E_{nn} . When superimposed on other orbitals with nodes, they leave the disposition of the nodes essentially unchanged. Following the procedure outlined in § 2, we leave them unaltered in the transformation to equivalent orbitals. We shall discuss CH_4 in some detail, the treatment of the inner shells of the other molecules being very similar.

Methane has tetrahedral symmetry T_d , so that a set of four equivalent orbitals can be formed from an A_1 orbital and a triply degenerate set T_2 (II, p. 22). Details of the group theory notation are given by Mulliken (1933) and Lennard-Jones (1934). The five occupied molecular orbitals are:

		symmetry
the carbon inner shell	ω_1	A_1
another completely symmetric orbital	ω_2	A_1
a triply degenerate set	$\omega_3, \omega_4, \omega_5$	T_2

The orbital ω_1 , corresponding to the inner shell, is left unaltered in the transformation from molecular orbitals to equivalent orbitals, the complete matrix for which is then

$$(t) = \frac{1}{2} \begin{pmatrix} 2 & 0 & 0 & 0 & 0 \\ 0 & 1 & 1 & 1 & 1 \\ 0 & 1 & 1 & -1 & -1 \\ 0 & 1 & -1 & 1 & -1 \\ 0 & 1 & -1 & -1 & 1 \end{pmatrix}. \quad (5.01)$$

If the equivalent sets are denoted by χ_1 and $\chi_2, \chi_3, \chi_4, \chi_5$ the corresponding energy matrix e_{mn} has values

$$\left. \begin{aligned} e_{11} &= E_{11}, & e_{m1} &= e_{1m} = 0 & (m \neq 1), \\ e_{mm} &= \frac{1}{4}(E_{22} + 3E_{33}) & (m \neq 1), \\ e_{mn} &= \frac{1}{4}(E_{22} - E_{33}) & (m, n > 1, m \neq n). \end{aligned} \right\} \quad (5.02)$$

The equations (2.01) are five in number, of which the last four are of similar form. The equation for the inner shell (χ_1) is

$$\{H + v'_1(x) - E_{11}\} \chi_1 = \sum_{m=2}^5 g_{m1}(x) \chi_m \quad (5.03)$$

or

$$\{H + g_{11}(x) - E_{11}\} \chi_1 = \sum_{m=2}^5 g_{m1}(x) \chi_m - 2 \sum_{m=2}^5 g_{mm}(x) \chi_1. \quad (5.04)$$

The first sets of terms on the right-hand side concerns the overlap between the inner shells and the outer orbitals. As the inner shell is concentrated near the carbon nucleus where the others are small, these may be neglected to a good approximation. The second set of terms on the right-hand side refer to the electrostatic effect of the outer orbitals on the inner one. If the outer orbitals (molecular or equivalent) form a spherically symmetric distribution which completely encloses the inner one, the terms $\sum_{m=2}^5 g_{mm}(x)$ give a constant potential energy in the significant region by classical electrostatic theory and do not alter the function χ_1 . These conditions are approximately satisfied in this molecule, so that we may neglect the effect of the second set of terms.

Similar considerations apply to the terms due to the attraction of the hydrogen nuclei, which are included in the Hamiltonian H . Even in molecules other than CH_4 , in which the other nuclei do not approximate so closely to a spherically symmetric distribution, they will be effectively screened by the electrons surrounding them and will only slightly perturb the inner shell.

With these approximations the equation (5.04) reduces to that for the ion C^{4+} . A better approximation would be to substitute on the right-hand side of the equation the corresponding functions for a neutral carbon atom. The calculations of Morse, Young & Hurwitz (1935) on analytical atomic wave functions, later modified by Coulson & Duncanson (1944), confirm that there is little difference between the K electrons of the ion C^{4+} and the core of neutral carbon.

To this degree of approximation the exchange terms with the inner shell can be neglected in the remaining four equations, which then become

$$\{H + v'_2(x) - \frac{1}{4}(E_{22} + 3E_{33})\} \chi_2 = \sum_{m=3}^5 \{g_{m2}(x) + \frac{1}{4}(E_{22} - E_{33})\} \chi_m, \quad (5.05)$$

and three others of similar form.

The inner shell affects (5.05) only through the condition of orthogonality and part of the term $v'_2(x)$. Since the potential outside a spherical distribution of charge is equal to that of the same charge concentrated at the centre, a reasonable approximation in dealing with this is to use the atomic inner shell function, restrict the other orbitals to be orthogonal to this, to prevent them from 'sagging' into regions of lower energy and then, having imposed this restriction, contract the core electronic charge into the C-nucleus.

If the central atom belongs to the second row of the periodic table, the L -shell electrons can be treated in a somewhat similar manner. In this case there is the further possibility of forming equivalent sets within the inner shell. For example,

the L -shell of silicon in SiH_4 can be regarded as made up of four equivalent orbitals along the four bond directions, but the orbitals will be concentrated near the silicon nucleus and will make little contribution to the bonding of the H-atoms.

(b) *Inner shells of general molecules*

In general, the inner shell of a particular atom in a molecule, such as Cl in CCl_4 , cannot be represented directly by a single molecular orbital. It is, however, always a member of a suitably chosen equivalent set. If we ignore the distorting effects of the other nuclei and electrons, such equivalent sets of inner shells can be obtained by transformations of molecular orbitals. In cases of ambiguity they can be fixed by the criterion of localization.

The equations for the inner shell equivalent sets now determine functions localized in the immediate neighbourhood of individual nuclei. As in the case of CH_4 they can be approximated to by means of atomic functions, so that by the reverse transformation we see that the corresponding molecular orbitals are approximately linear combinations of atomic orbitals.

A simple example is ethane C_2H_6 , in which the molecular orbitals are bonding and antibonding, symmetrical and antisymmetrical respectively in the plane bisecting the C-C bond, while the individual inner shells form a set of two equivalent orbitals. Another example is carbon tetrachloride CCl_4 , in which the carbon inner shell forms an equivalent set on its own and the chlorine shells form a set of four equivalent orbitals. The molecular orbitals corresponding to the latter are linear combinations belonging to representations A_1 and T_2 of the tetrahedral symmetry group T_d .

6. THE SIGNIFICANCE OF ELECTRONS ASSIGNED TO LONE PAIRS

Lone pairs of electrons, that is, pairs of electrons with opposite spins in orbitals in the region of one nucleus only, and not shared by two, can also be discussed in terms of molecular or equivalent orbitals. As in the case of inner shells, we shall deal first with molecules of which only one atom has lone pairs, using the central hydrides ammonia and water as examples. Other molecules of this type can be discussed in a similar way.

(a) *Lone pairs in ammonia*

The molecule NH_3 has trigonal symmetry C_{3v} , so that a set of three equivalent orbitals can be formed by a combination of an A_1 orbital and a degenerate pair with symmetry E (II, p. 21). Now the occupied molecular orbitals are:

		symmetry
the nitrogen inner shell (K electrons)	ω_1	A_1
two further completely symmetric orbitals	ω_2, ω_3	A_1
the degenerate pair	ω_4, ω_5	E

The inner shell function ω_1 , as in CH_4 , is concentrated close to the nitrogen nucleus and takes no part in the chemical bonding. We shall suppose it dealt with by the method discussed in § 5. ω_2 and ω_3 are molecular orbitals defined uniquely by the

condition $E_{23} = 0$. Whether or not one of them can be used directly in the formation of suitable equivalent orbitals depends on the regions of maximum absolute value and the position of the nodal surfaces. We can say that if one of them, say ω_2 , is concentrated mainly on the hydrogen side of the nitrogen nucleus, it may be called the bond molecular orbital, while ω_3 will be on the remote side and will correspond to the lone pair. Evidently ω_2 will have the larger ionization potential, as it is closer to the three hydrogen nuclei. Under these circumstances ω_2 is obviously the best A_1 molecular orbital to use in the formation of the three equivalent bond orbitals, as it overlaps more with other orbitals in the N-H regions than does ω_3 .

The energies and distributions of the orbitals ω_2 and ω_3 can be illustrated by reference to an atom in a strong electric field. In an isolated atom the sequence of orbitals is $1s, 2s, 2px, 2py, 2pz$, the last three being degenerate, so that any set of orthogonal directions x, y and z is equally valid. In a strong electric field, one direction is prescribed, say the x direction, and the orbitals must now be described by reference to their symmetry about this direction. The orbitals $2s$ and $2px$ change their forms and their energies, but each retains axial symmetry about the field. One may be described as $2\sigma_1$ and the other as $2\sigma_2$. (These correspond to the orbitals labelled ω_2 and ω_3 in the ammonia molecule.) One of them, say $2\sigma_1$ takes up a distribution with a well defined region of maximum probability in the direction of the electric field, while the other, $2\sigma_2$, has a region of maximum probability on the remote side. This feature of $2\sigma_2$ is due to the fact that it must remain orthogonal to $2\sigma_1$. This is analogous to the discussion of equivalent orbitals earlier in the paper. This behaviour of $2\sigma_1$ and $2\sigma_2$ can be simulated by taking linear combinations of $2s$ and $2px$ and using perturbation theory.

We now apply the transformation obtained in II (equation (4.04)) to ω_2, ω_4 and ω_5 , leaving the others unaltered. The equations for the bond equivalent orbitals χ_1, χ_2 and χ_3 now have the form

$$\begin{aligned} \{H + v'_1(x) - \frac{1}{3}(E_{22} + 2E_{44})\} \chi_1 \\ = \{g_{21}(x) + \frac{1}{3}(E_{22} - E_{44})\} \chi_2 + \{g_{31}(x) + \frac{1}{3}(E_{22} - E_{44})\} \chi_3 + g_{41}(x) \chi_4, \end{aligned} \quad (6.01)$$

and two similar equations, while that for the lone pair χ_4 is

$$\{H + v'_4(x) - E_{44}\} \chi_4 = g_{14}(x) \chi_1 + g_{24}(x) \chi_2 + g_{34}(x) \chi_3. \quad (6.02)$$

Thus the transformation has reduced the problem to the determination of two functions only.

If the molecular orbitals ω_2 and ω_3 do not give satisfactory localized functions when transformed in the above way, an alternative set can be obtained from a symmetry orbital formed by a normalized linear combination

$$s_2 = \omega_2 \cos \lambda + \omega_3 \sin \lambda. \quad (6.03)$$

The remaining orbital would then be

$$s_3 = -\omega_2 \sin \lambda + \omega_3 \cos \lambda. \quad (6.04)$$

The equations (6.01) and (6.02) have to be replaced by the more general forms

$$\left. \begin{aligned} \{H + v'_1(x) - e_{11}\} \chi_1 &= \sum_{n=2}^4 \{g_{n1}(x) - e_{n1}\} \chi_n, \\ \{H + v'_4(x) - e_{44}\} \chi_4 &= \sum_{n=1}^3 \{g_{n4}(x) - e_{n4}\} \chi_n. \end{aligned} \right\} \quad (6.05)$$

The solutions of (6.05) contain λ as an arbitrary parameter which can be adjusted to give maximum localization.

(b) *Lone pairs in the water molecule*

The five occupied molecular orbitals for the molecule H_2O are (Mulliken 1933):

		symmetry
the oxygen inner shell	ω_1	A_1
one B_1 orbital with node in plane midway between OH bonds	ω_2	B_1
two orbitals with the complete symmetry of the nuclear framework	ω_3, ω_4	A_1
an orbital with a node in the nuclear plane	ω_5	B_2

Of these, ω_5 will have the smallest ionization potential, as these electrons will be kept away from the nuclear plane by the node. Of ω_3 and ω_4 , the one with the larger ionization potential will be concentrated more on the side of the hydrogen atoms. Suppose this is ω_3 . Then if this is used in conjunction with ω_2 to form equivalent bond orbitals (I, equation (5.02)), ω_4 and ω_5 are left as lone pair molecular orbitals. In this case there is the further possibility of applying a similar transformation to these to get two equivalent lone pairs, mainly concentrated in regions above and below the nuclear plane.

If the bond equivalent orbitals are χ_1 and χ_2 and the others χ_3, χ_4 , the equations determining them are

$$\left. \begin{aligned} \{H + v'_1(x) - \tfrac{1}{2}(E_{22} + E_{23})\} \chi_1 &= \{g_{21}(x) + \tfrac{1}{2}(E_{33} - E_{22})\} \chi_2 + g_{31}(x) \chi_3 + g_{41}(x) \chi_4, \\ \{H + v'_3(x) - \tfrac{1}{2}(E_{44} + E_{55})\} \chi_3 &= g_{13}(x) \chi_1 + g_{23}(x) \chi_2 + \{g_{43}(x) + \tfrac{1}{2}(E_{44} - E_{55})\} \chi_4. \end{aligned} \right\} \quad (6.06)$$

and two similar equations for χ_2 and χ_4 .

As in the case of NH_3 the two sets of equivalent orbitals may be obtained alternatively from symmetry orbitals formed from ω_3 and ω_4 , in which case equations (6.06) must be replaced by the more general ones corresponding to (6.05). A fuller discussion of these transformations and the dependence of the functions on the nuclear configuration will be given in part V.

(c) *Lone pairs in larger molecules*

In larger molecules where two or more atoms possess lone pairs, the corresponding molecular orbitals may be non-localized, when it will be necessary to take linear combinations to obtain functions concentrated near one nucleus. If localization to atoms leave two lone pair orbitals near certain nuclei, other linear combinations may localize them further.

For example, in hydrogen peroxide H_2O_2 , lone pair orbitals will be formed at each O atom by adding and subtracting molecular orbitals spread over both oxygen nuclei. It may then be possible to apply a further transformation similar to that discussed for H_2O and NH_3 above.

7. FACTORS WHICH DETERMINE THE SPATIAL ARRANGEMENT OF ATOMS IN A MOLECULE

In § 3 we give a general expression for the energy of a molecule. This was divided for convenience into four parts. One gave the contribution made by the electrostatic repulsion of the nuclei. Another part represented the kinetic energy of the electrons and their potential energy due to the nuclei alone. This part depends on the form of the molecular orbitals to which the electrons are assigned, but not on a transformation of these orbitals, provided this transformation is of the type discussed in this paper. The kinetic energy of the electrons has the same value whether they are regarded as in molecular orbitals or equivalent orbitals. The same is true of the potential energy of the electrons due to the nuclei alone.

The other parts of the energy arise from the interaction of the electrons with each other. These include the electrostatic repulsion of electrons in their various charge distributions, a contribution which is readily understood, and the 'exchange' interaction, whose meaning is not so easily apprehended nor its influence on molecular structure so easily appreciated. The object of this investigation is to choose such orbitals for the electrons as will make this exchange part as small as possible. This is done by using equivalent orbitals, and some of these may be interpreted as localized bonds, and others as directed orbitals in particular atoms, usually described as lone pairs. The remaining occupied orbitals will be more tightly bound than these and will constitute the inner shells.

If the orbitals representing bonds are denoted by the symbol b , those representing the lone pairs by l , and those representing inner orbitals by i , the electrostatic energy, corresponding to $\mathfrak{J}$ and $\mathfrak{K}$ given in equation (3.07), consists of the following parts

$$\mathfrak{J} = \mathfrak{J}_{ii} + \mathfrak{J}_{bb} + \mathfrak{J}_l + \mathfrak{J}_{ib} + \mathfrak{J}_{il} + \mathfrak{J}_{lb}, \quad (7.01)$$

$$-\mathfrak{K} = -\mathfrak{K}_{ii} - \mathfrak{K}_{bb} - \mathfrak{K}_l - \mathfrak{K}_{ib} - \mathfrak{K}_{il} - \mathfrak{K}_{lb}. \quad (7.02)$$

In a similar way we can separate the contributions of the bond, lone pair and inner shell orbitals to the kinetic energy and potential energy of the electrons in the field of the nuclei.

The terms involving the inner shells will be important parts of the absolute energy of the molecule, but will not depend very much on the disposition of the bonds. They will also partly counteract the field of the nuclei and so screen them that outer electrons will be distributed as though the nuclear charges were effectively reduced.

The shape of the molecule and its external properties will be determined largely by the forces which do not involve the inner electrons. The electron repulsions which are effective in this respect are represented by the terms of the following expression:

$$\mathfrak{J}_{bb} + \mathfrak{J}_l + \mathfrak{J}_{bl} - \mathfrak{K}_{bb} - \mathfrak{K}_l - \mathfrak{K}_{bl}. \quad (7.03)$$

If we cancel the common parts of terms such as $\mathfrak{J}_{bb}$ and $\mathfrak{R}_{bb}$, this may be written

$$j_{bb} + j_{ll} + j_{bl} - k_{bb} - k_{ll} - k_{bl}, \quad (7.04)$$

where j_{bb} represents the Coulomb interaction between bonds and $-k_{bb}$ the corresponding exchange interaction. In addition, there are the corresponding terms in the kinetic energy and electron-nuclear potential energy involving bond and lone pair electrons. Finally, the internuclear repulsions are an important factor in determining the size and shape of the molecule.

The point on which we wish to lay emphasis here is the occurrence within formulae (7.03) and (7.04) of terms involving the lone pairs. The distribution in space of these lone pairs will depend on the structure of the molecule. A change in the bond angles will produce a change of the relative orientation of the lone pairs. Thus, through their electrostatic interaction, lone pairs may resist alteration in bond angles and consequently be important factors in determining the equilibrium configurations of molecules. This function of lone pairs does not seem to have been recognized hitherto. It will be investigated in a subsequent paper and shown to be an important factor in determining the structure of molecules such as water.

One of the authors (J.A.P.) is indebted to the Department of Scientific and Industrial Research for a grant.

REFERENCES

- Coulson, C. A. & Duncanson, W. E. 1944 *Proc. Roy. Soc. Edinb.* **62 A**, 37.
 Hall, G. G. & Lennard-Jones, Sir J. 1950 *Proc. Roy. Soc. A*, **202**, 155 (part III).
 Lennard-Jones, Sir J. 1934 *Trans. Faraday Soc.* **30**, 70.
 Lennard-Jones, Sir J. 1949 *Proc. Roy. Soc. A*, **198**, 1 (part I).
 Lennard-Jones, Sir J. 1949 *Proc. Roy. Soc. A*, **198**, 14 (part II).
 Morse, P. M., Young, L. A. & Hurwitz, E. S. 1935 *Phys. Rev.* **48**, 948.
 Mulliken, R. S. 1933 *Phys. Rev.* **43**, 297.
 Slater, J. C. 1930 *Phys. Rev.* **36**, 57.

The kinetics of the interaction of atomic hydrogen with olefines

V. Results obtained for a further series of compounds

BY H. W. MELVILLE, F.R.S. AND J. C. ROBB

Chemistry Department, University of Birmingham

(Received 18 January 1950)

In this paper the efficiency of interaction of a hydrogen atom with a series of olefines has been determined, the olefines being members of the series obtained by progressively replacing the hydrogen atoms of ethylene by methyl radicals. The interesting generalization which emerges from this is that the efficiency of interaction does not vary very much with the nature of the alkyl substituents in the molecule, and calculations involving the heats of addition of a hydrogen atom to a double bond confirm this generalization. The data presented here are discussed critically in relation to information available on the reaction of CCl_3 radicals with olefines and of alkyl radicals with olefines, complete general agreement being demonstrated.

INTRODUCTION

The previous parts of this series (Melville & Robb 1949) have illustrated a new technique for measuring the absolute efficiency of fast reactions of atomic hydrogen. The method described consisted in essence of having a surface of molybdenum trioxide so situated within a reaction vessel as to compete with an olefine molecule for the hydrogen atoms produced by the mercury photo-sensitized decomposition of molecular hydrogen. The present paper is an extension of the results already obtained by this method. A series of olefines has now been studied completely, whereby the hydrogen atoms of ethylene have been progressively replaced by methyl groups. In addition, some investigation has been made with reference to the interaction of a hydrogen atom with hexane, benzene, butadiene, acetylene, and oxygen.

APPARATUS

The apparatus used is basically the same as that already described, but with a few alterations to improve the precision of the method.

Colorimeter

In the former model, light was shone on the molybdenum oxide surface by means of a prism, but this has been replaced by a direct beam produced by a 12 V, 2.2 W bulb, under-run at 10 V and focused by a small lens to give an approximately parallel beam of light on the oxide surface (figure 1). The under-running of the lamp has a two-fold purpose. First, it prevents any geometrical change in the lamp filament arrangement, and secondly, it reduces any possible thermal effect on the photocell.

Reaction system

With some of the higher hydrocarbons (most noticeably with trimethyl ethylene), solution of the olefine in the greases and waxes became a problem. Accordingly, all important joints were sealed with mercury. These seals were so successful that they were retained even when not essential in view of the increased simplicity in handling the system. The main seal is that between the flanged top of the reaction vessel and the silica window. The method adopted was to construct a ring of polythene which fitted round the flanged top and formed a channel into which mercury could be poured (figure 2). This ring was sealed to the flanged top by running paraffin wax into the small spaces between the vessel and the polythene ring and allowing it to set hard. It was essential that no mercury vapour should escape from the seal and saturate the air above the reaction vessel, since this would have caused premature

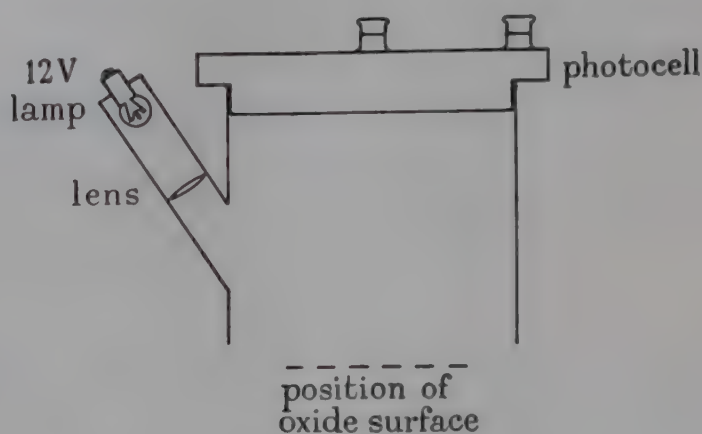


FIGURE 1. Colorimeter.

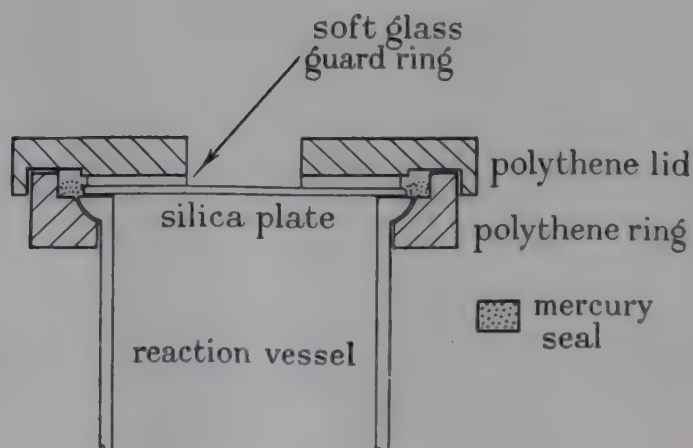


FIGURE 2. Detail of mercury-sealed flange joint.

absorption of 2537 Å radiation. A polythene cover was made to fit closely over the silica window and guard ring of soft glass and thus prevent vaporization of mercury in the seal. The surfaces of the polythene cover were lightly greased to ensure that the fit was perfect. No trouble was experienced at any time owing to mercury vapour leaking out from the seal. When the reaction vessel was evacuated,

mercury ran from the seal between the ground faces of the joint to a depth of about 1 mm. and formed a perfectly tight seal.

The bottom joint of the reaction vessel was sealed in the conventional manner by fitting a glass sleeve over the joint and filling the annular space with mercury. This joint, of course, could not be turned once the vessel was evacuated, but this was no disadvantage since the movable column could be preset before evacuation of the vessel. No trouble was experienced owing to evaporation of mercury from this surface. The side joint, connecting the reaction vessel to the vacuum system, was sealed by fitting a glass sleeve over it, closed at both ends by small sections of rubber tube. In the centre of the sleeve a hole was blown out with raised edges which were ground flat. Mercury was introduced through this hole to seal the joint and the hole closed by a small waxed ground-glass plate.

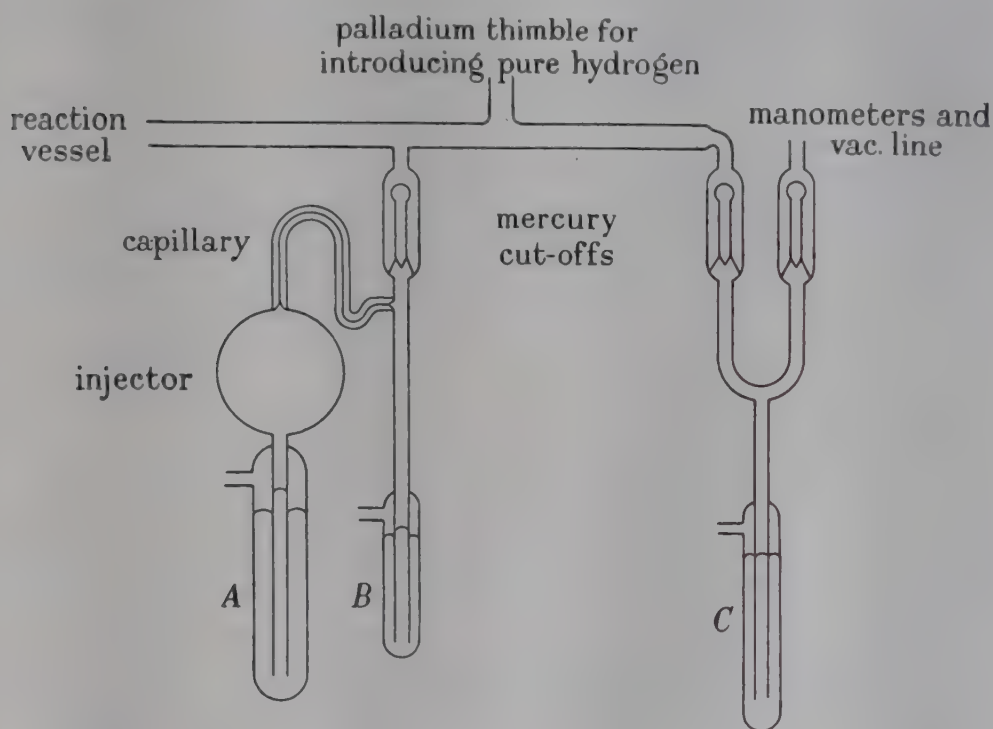


FIGURE 3. Modification of apparatus to remove necessity for taps.

The injector, by means of which a measured amount of olefine was introduced had also to be controlled by mercury cut-offs instead of taps (figure 3). The method of operation of the injector was to introduce olefine to the system at the required pressure with the mercury in reservoirs *A*, *B* and *C* lowered. Cut-off *B* was then raised until mercury just flowed round the capillary bend *D*, thus sealing the injector. The rest of the system was then pumped out, cut-off *C* raised and hydrogen introduced by electrically heating the palladium thimble until the required pressure of hydrogen diffused through. When the olefine is required in the reaction vessel, it is necessary only to raise the mercury in *A*, thus compressing the olefine and pushing the sealing mercury out of the capillary. The mercury in the injector is then raised to expel all the olefine into the reaction vessel. When the mercury is dropped in *A*, the capillary becomes automatically sealed again with mercury.

Control of light duration

Previously, manual operation of a shutter was employed, but this method has been replaced by an automatic shutter, giving accurately equal exposures of about 30 sec. duration. The circuit is shown in figure 4. *A* is a Tufnol disk, fitted over the minute spindle of an a.c. clock motor, and has two studs fitted diametrically opposite one another and ground flush with the Tufnol disk. To the minute spindle of the motor is fitted a phosphor-bronze strip *B*, having a small ball-bearing let into a small hole in the strip near the end. The ball is so placed as to make contact with the studs when it sweeps round the disk. When contact is made, current passes through the relay *S*, breaking *S*₁ and stopping the motor, opening *S*₂ and permitting the electromagnetically operated flap type of shutter to close by means of return springs. When an exposure is required, press-button *K* is pressed until the motor drives the contractor strip off the stud which then releases the relay, closes *S*₁ to maintain the motor drive and closes *S*₂ to open the shutter. *K* can be released and the light exposure ceases when the other stud is contacted.

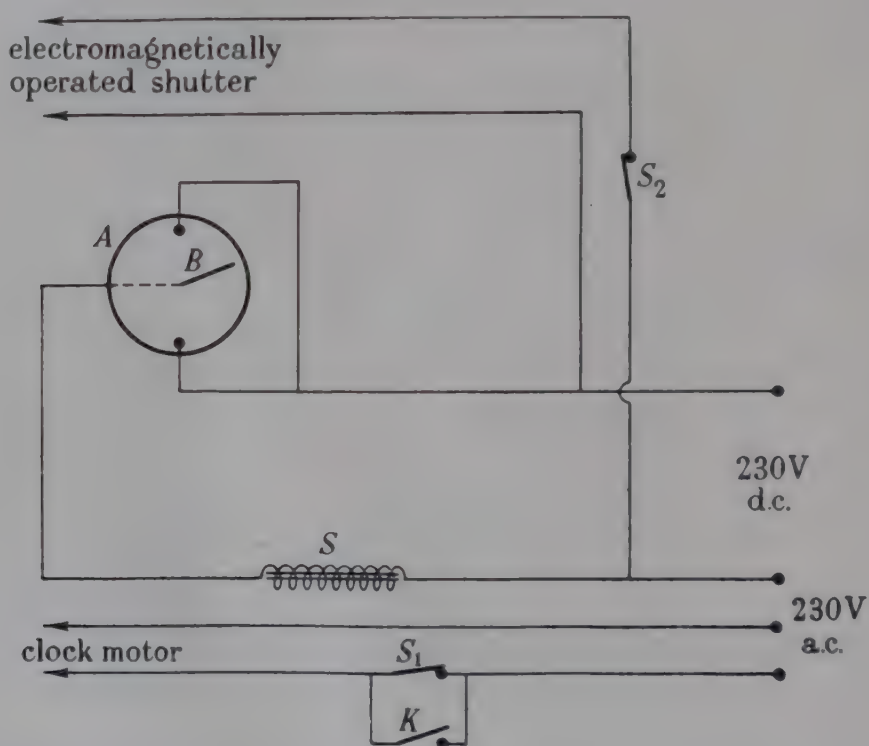


FIGURE 4. Circuit of shutter timing device.

Photocell circuit measurement

The bridge for measuring the photocell circuit is run off a d.c. source stabilized to 1 part in 10,000 to eliminate the need for accumulators. The mercury lamp and the lamp in the colorimeter are both run off a.c., electronically stabilized to 1 part in 10,000.

Techniques

In the previous parts all resistance measurements were converted and expressed in terms of current. It is now considered to be much less laborious to work in terms of resistance throughout.

The relationship deduced previously was

$$kt + k_1 = 1/(C' - k_2).$$

Since $C = E/R$, where E is the voltage applied to the galvanometer and series variable resistance, and k , k_1 and k_2 are constants, then

$$kt + k_1 = 1/(E/R' - k_2),$$

$$kt + k_1 = R'/(E - k_2 R'),$$

$$E(kt + k_1) = R' \left/ \left(E - \frac{k_2}{E} R' \right) \right.$$

Replacing by new constants,

$$k_6 t + k_7 = R'/(1 - k_8 R').$$

Thus by determining the value of k_8 and plotting the right-hand side against t , a straight line will be obtained as before, whose slope will be proportional to the rate of addition of hydrogen atoms to the molybdenum oxide surface.

The previous method of conducting a run consisted in measuring the blueing due to atomic hydrogen alone for 2.5 min., then with hydrogen and olefine for a further 2.5 min. and finally with pure hydrogen again for 2.5 min. The value of k_2 was computed by taking two values for current at t_1 and t_2 on the first portion of the run and two values for t_3 and t_4 from the last portion. By the same method, for resistances R_1 and R_2 at t_1 and t_2 ,

$$k_6 t_1 + k_7 = R_1/(1 - k_8 R_1), \quad (1)$$

$$k_6 t_2 + k_7 = R_2/(1 - k_8 R_2). \quad (2)$$

Eliminating k_7 ,

$$k_6(t_2 - t_1) = \frac{R_2}{1 - k_8 R_2} - \frac{R_1}{1 - k_8 R_1}. \quad (3)$$

Similarly for t_3 and t_4

$$k_6(t_4 - t_3) = \frac{R_4}{1 - k_8 R_4} - \frac{R_3}{1 - k_8 R_3}. \quad (4)$$

Eliminating k_6 ,

$$\frac{t_2 - t_1}{t_4 - t_3} = \frac{\frac{R_2}{1 - k_8 R_2} - \frac{R_1}{1 - k_8 R_1}}{\frac{R_4}{1 - k_8 R_4} - \frac{R_3}{1 - k_8 R_3}},$$

from which k_8 can be determined.

It has, however, proved satisfactory to conduct the run in a rather different way. A run is done first with hydrogen alone for 5 min., then olefine is injected and the run continued for a further 3 min. The value of k_8 can be calculated by taking the resistance values at $t_1 = 0$, $t_2 = 2.5$ min., $t_3 = 5.0$ min. Then, as above,

$$k_6(t_2 - t_1) = \frac{R_2}{1 - k_8 R_2} - \frac{R_1}{1 - k_8 R_1},$$

$$k_6(t_3 - t_2) = \frac{R_3}{1 - k_8 R_3} - \frac{R_2}{1 - k_8 R_2}.$$

Then, since $t_2 - t_1 = t_3 - t_2$,

$$\frac{R_2}{1 - k_8 R_2} - \frac{R_1}{1 - k_8 R_1} = \frac{R_3}{1 - k_8 R_3} - \frac{R_2}{1 - k_8 R_2},$$

from which k_8 can be readily determined.

EXPERIMENTAL PROCEDURE

The most satisfactory method of conducting an experimental run was as follows. The olefine was introduced as already described, mercury vaporized, and after introduction of the hydrogen, two light exposures, each of half a minute duration, were given and then the resistance in the photocell measuring unit adjusted and recorded after every subsequent half-minute exposure for a total of 5 min. The olefine was then introduced and the system left for 3 min. to establish a uniform concentration of olefine in the reaction system. Resistance readings were then recorded every half-minute for a further 3 min. After calculation of k_8 , the values of $\frac{R}{1 - k_8 R}$ were obtained and plotted against time. From the slopes of the two parts of the run, the ratio was obtained. The ratios for spacings of 2.5 and 1.25 mm. were then set on the calculator described in part III, and the value of 'b' determined. From this, as previously, the velocity constant, k , and collision efficiency were determined.

As before, mercury was vaporized in the reaction vessel before starting each run.

Purification of hydrocarbons

The hydrocarbons used in this investigation were purified by conventional methods. These were degassed by condensation in liquid air, pumping-off non-condensable gas and subsequent melting, this procedure being repeated until no dissolved gas remained.

EXPERIMENTAL RESULTS

Trimethyl ethylene

Figure 5 shows runs 32, 33 and 34, conducted with trimethyl ethylene at a spacing of 0.25 mm. on molybdenum oxide. The value of k_8 for these runs was 0.044, 0.035, 0.033, respectively, and the ratio of the slopes of the portions with and without olefine, which was present from $t = 2.5$ min. to $t = 5.0$ min., were respectively 1.72, 1.74, 1.70, giving an average slope of 1.72.

Figure 6 shows runs 39, 40 and 41, with the same olefine but at a spacing of 1.25 mm. on molybdenum oxide. k_8 for these runs was respectively 0.048, 0.0455, 0.0497, and the ratio of slopes 1.46, 1.48, 1.39, giving an average 1.44.

From these values for slopes at spacing 1.25 and 2.5 mm. the value of b obtained by means of the calculator described in part III is 14 and the blueing coefficient on MoO_3 is 1.89. The molecular diameter obtained by measurement of the appropriate Stuart model is 7.5×10^{-8} cm. This gives for the value of D , the diffusion coefficient

$$D = 8 \times 10^4 / (102.7 p_2 + 290 p_3),$$

where p_2 is the pressure of hydrogen in the vessel and p_3 the pressure of olefine. Substituting the values for p_2 and p_3 gives $D = 81 \text{ cm.}^2/\text{sec.}$ By definition

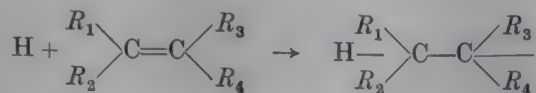
$$b = \sqrt{(g/D)},$$

$$g = 1.6 \times 10^4.$$

Further

$$g = kn,$$

where k is the velocity constant of the reaction



and n is the number of olefine molecules per cm.^3

$$k = \frac{1.6 \times 10^4}{0.4 \times 3.6 \times 10^{16}} \\ = 1.1 \times 10^{-12} \text{ cm.}^3 \text{ molecules}^{-1} \text{ sec.}^{-1},$$

the collision number, computed in the usual way is 1.9×10^{-9} ,

$$\text{Collision efficiency} = \frac{1.1 \times 10^{-12}}{1.9 \times 10^{-9}} \\ = 5.8 \times 10^{-4}.$$

The above investigation was carried out before the mercury sealed system was developed, and it was considered advisable to repeat the work using the improved

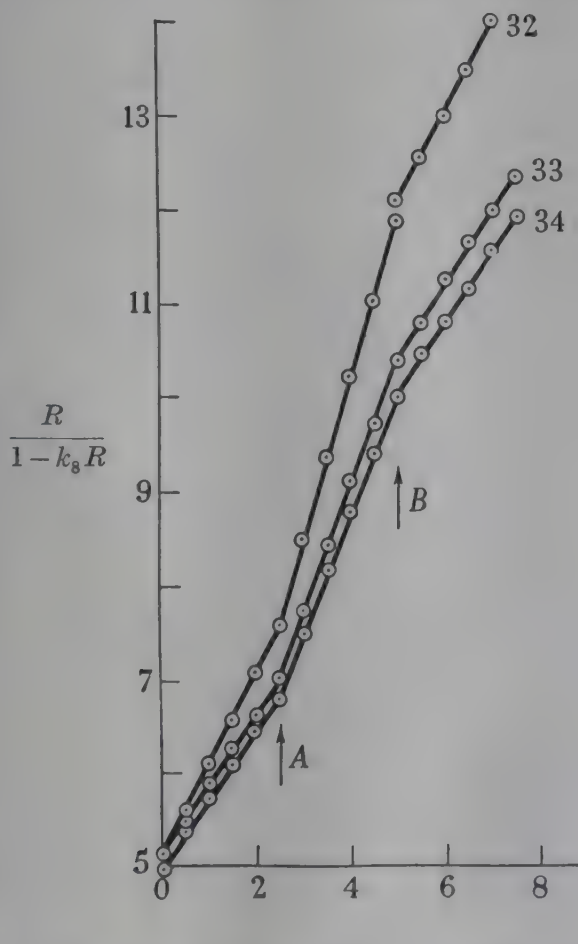


FIGURE 5. Trimethyl ethylene : spacing 2.5 mm., 8.4 mm. H_2 : 0.41 mm. olefine at A; 8.4 mm. H_2 at B.

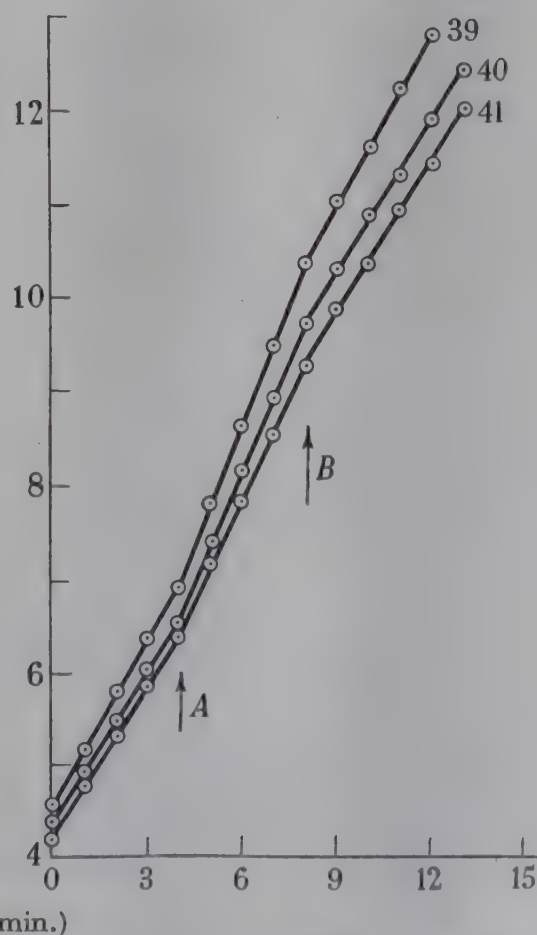


FIGURE 6. Trimethyl ethylene : spacing 1.25 mm., 8.4 mm. H_2 : 0.41 mm. olefine at A; 8.4 mm. H_2 at B.

reaction vessel arrangements. A further change in technique was made. The previous technique adopted for a spacing of 2.5 mm. was to conduct the run with hydrogen for a period of 2.5 min. then for a further period of 2.5 min. after injection of the olefine and finally for 2.5 min. with pure hydrogen. For the 1.25 mm. spacing, the time periods were 4, 4, and 5 min. respectively. In this new type of run the initial period with pure hydrogen was increased to 5 min. and the olefine period to 3 min. The final hydrogen period was omitted and the value of the constant k_s computed from the single hydrogen period. This technique gave equally good results and simplified the conduct of the run.

Figure 7 shows runs 81, 82 and 83 at a spacing of 2.5 mm. The value of k_s for these runs is 0.0625, 0.0655, 0.0635, and the ratio of slopes, 2.14, 2.26, 2.24, giving an average of 2.21.

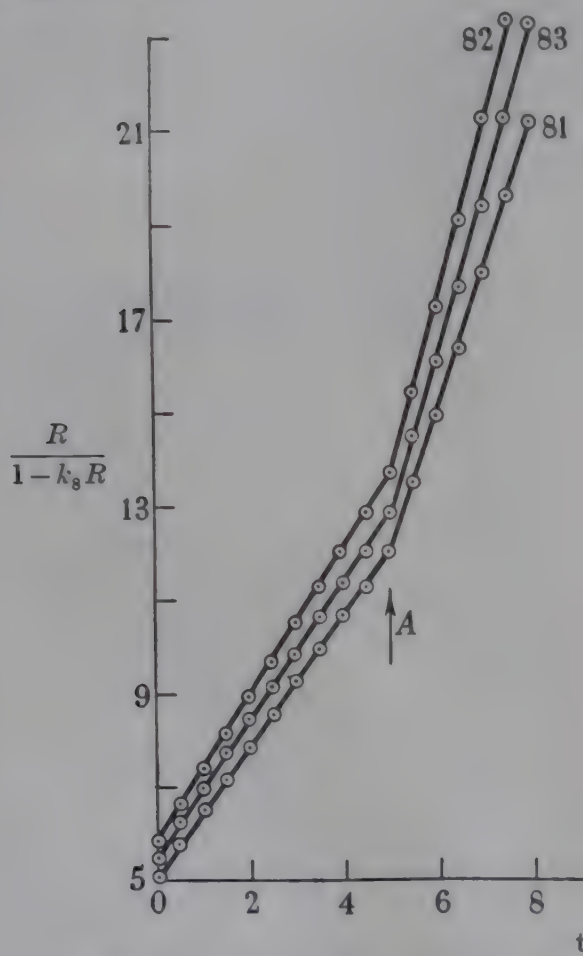


FIGURE 7. Trimethyl ethylene: spacing 2.5 mm., 8.3 mm. H_2 : 0.41 mm. olefine at A.

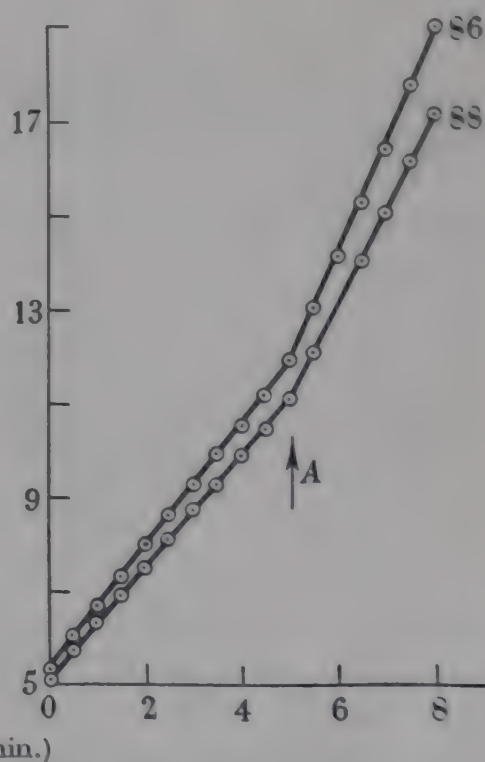


FIGURE 8. Trimethyl ethylene: spacing 1.25 mm., 8.3 mm. H_2 : 0.41 mm. olefine at A.

Figure 8 gives runs 86, 88 for a spacing of 1.25 mm. The values of k_s are 0.769, 0.720, and the ratios of slopes 1.73, 1.69, giving an average of 1.71. This leads to a value for b of 13 and a blueing coefficient of 2.55. As before, this leads to a value of

$$g = 1.4 \times 10^4,$$

whence

$$k = 9.5 \times 10^{-13} \text{ cm.}^3 \text{ molecules}^{-1} \text{ sec.}^{-1},$$

and also the collision efficiency is 5×10^{-4} .

The values obtained by this method agree satisfactorily with those already obtained by the other, the most striking variation being the change of blueing coefficient from 1.89 to 2.55. The explanation seems to be that the blueing rate for any radical is to a certain extent dependent on the previous history of that particular sample of molybdenum oxide. The former method takes no account of the fact that the blueing rate may be different for hydrogen after the surface acquires a number of alkyl radicals. The latter method, being independent of a hydrogenation rate subsequent to radical addition, should give a more exact result although the difference as seen above is so small as to be unimportant as far as this investigation is concerned.

Tetra-methyl ethylene

Figure 9 shows runs 45, 47, 48, 49 at a spacing of 2.5 mm. The values for k_8 are 0.031, 0.043, 0.036, 0.050, and the blueing ratios are 1.71, 1.67, 1.68, 1.62, giving an average value of 1.67.

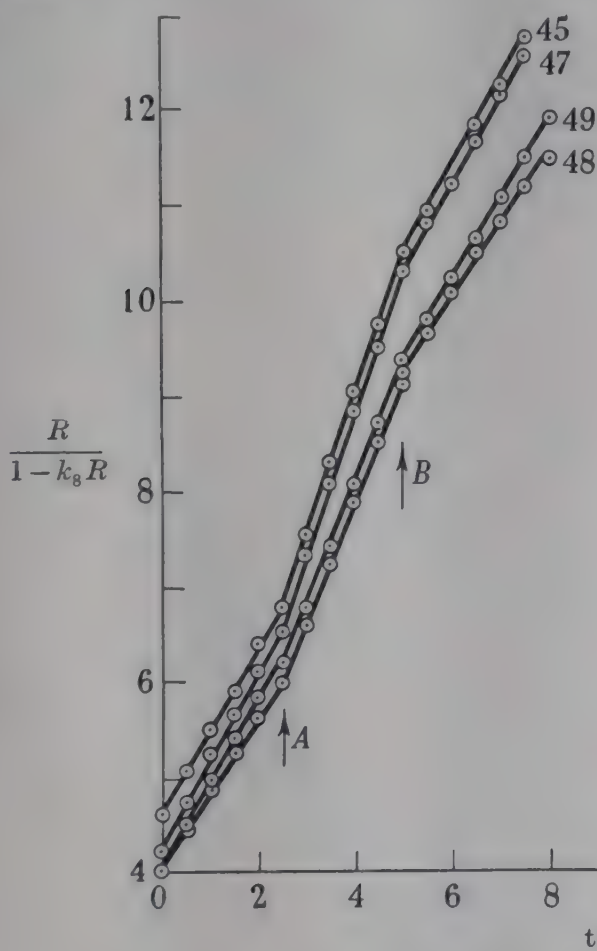


FIGURE 9. Tetra-methyl ethylene: spacing 2.5 mm., 8.1 mm. H_2 : 0.41 mm. olefine at A; 8.1 mm. H_2 at B.

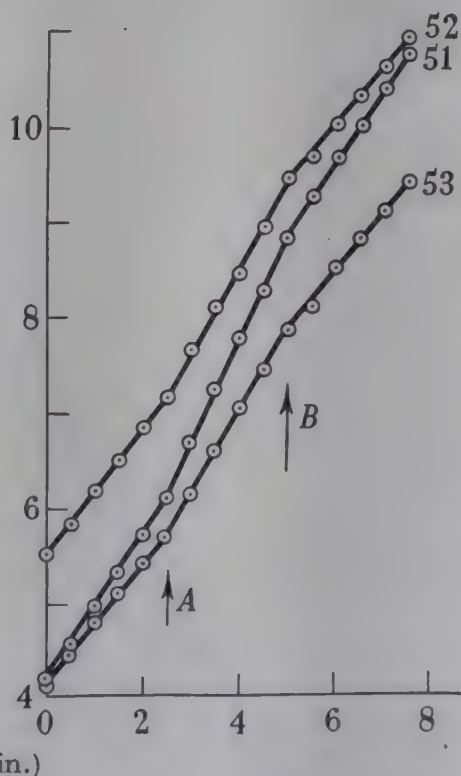


FIGURE 10. Tetra-methyl ethylene: spacing 1.25 mm., 8.1 mm. H_2 : 0.41 mm. olefine at A; 8.1 mm. H_2 at B.

Figure 10 shows runs 51, 52, 53 at a spacing of 1.25 mm. The value of k_8 are 0.0476, 0.0440, 0.0363, and the blueing ratios are 1.37, 1.35, 1.35, giving an average value of 1.36. These values for spacings 2.5 and 1.25 mm. give a value for b of 11 and the

blueing coefficient is 1.94. The diameter of the molecule by measurement is 7.5×10^{-8} cm., hence

$$\begin{aligned} D &= 8 \times 10^4 / (102.7 p_2 + 290 p_3) \\ &= 84, \\ g &= b^2 D = kn \\ &= 1.02 \times 10^4 = kn, \\ k &= 1.02 \times 10^4 / (0.41 \times 3.6 \times 10^{16}) \\ &= 6.9 \times 10^{-13} \text{ cm.}^3 \text{ molecules}^{-1} \text{ sec.}^{-1}. \end{aligned}$$

As before, the collision number is computed at 1.9×10^{-9} . Hence the collision efficiency

$$\begin{aligned} &= 6.9 \times 10^{-13} / 1.9 \times 10^{-9} \\ &= 3.6 \times 10^{-4}. \end{aligned}$$

Butene-2 (Trans)

Figure 11 shows runs 93, 94, 96, with molybdenum oxide at a spacing of 2.5 mm. The respective values of k_8 are 0.0645, 0.0617, 0.0589 and the blueing ratios are 1.91, 1.94, 1.99, giving an average value of 1.95.

Figure 12 shows runs 98, 99, 100. The values of k_8 are 0.0664, 0.0689, 0.0700 and the blueing ratios are 1.56, 1.57, 1.59, giving an average value of 1.57.

From these values, $b = 13$ and the blueing coefficient is 2.26.

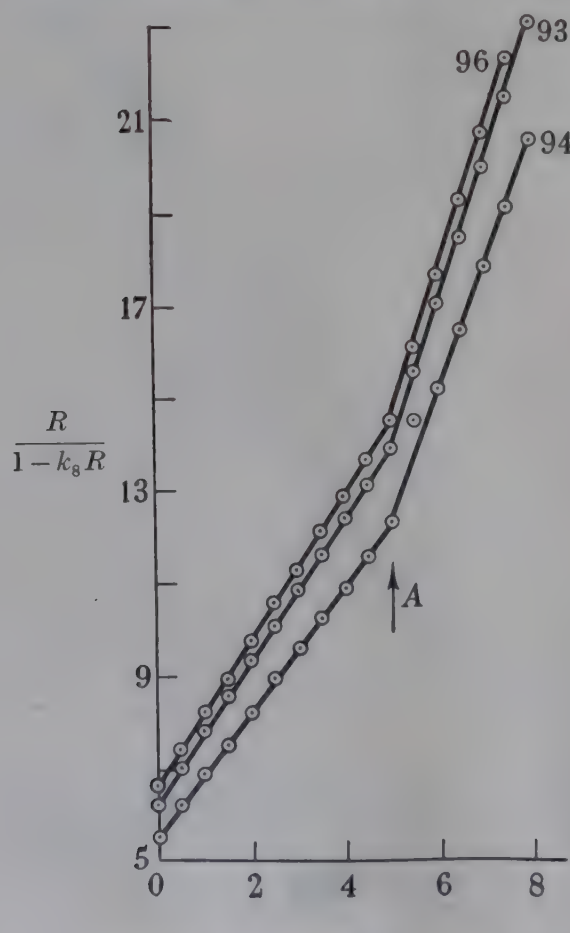


FIGURE 11. Butene-2 (*trans*): spacing 2.5 mm., 8.3 mm. H_2 : 0.41 mm. olefine at A.

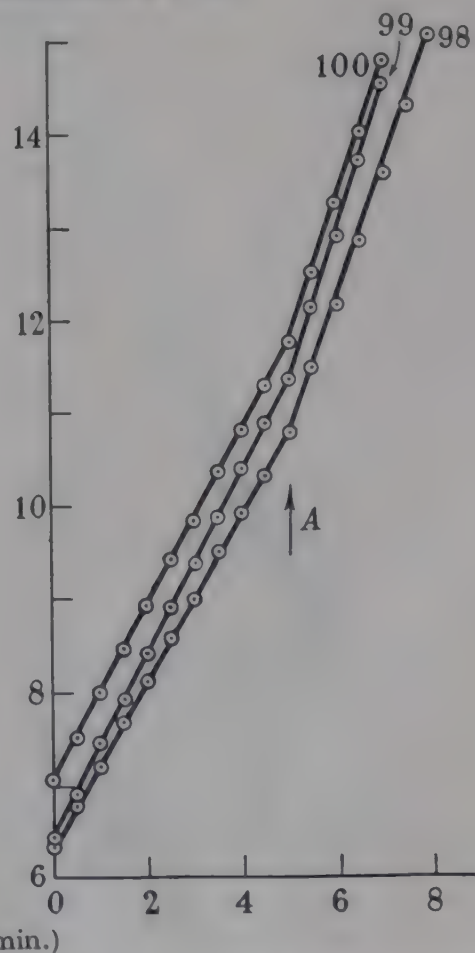


FIGURE 12. Butene-2 (*trans*): spacing 1.25 mm., 8.3 mm. H_2 : 0.41 mm. olefine at A.

From model measurements, the diameter of the molecule is found to be 7.5×10^{-8} cm., and from this value

$$\begin{aligned} D &= 8 \times 10^4 / (102.7 p_2 + 290 p_3) \\ &= 82.5, \\ g &= b^2 D = 1.39 \times 10^4, \\ k &= g/n = 1.39 \times 10^4 / 0.41 \times 3.6 \times 10^{16} \\ &= 9.5 \times 10^{-13} \text{ cm.}^3 \text{ molecules}^{-1} \text{ sec.}^{-1}. \end{aligned}$$

In the usual way, the collision number is computed at 1.9×10^{-9} ,

$$\begin{aligned} \text{Collision efficiency} &= 9.5 \times 10^{-13} / 1.9 \times 10^{-9} \\ &= 5 \times 10^{-4}. \end{aligned}$$

Butene-2 (cis)

Figure 13 shows runs 111, 112, 113, 114 for molybdenum oxide at a spacing of 2.5 mm. The values of k_s are 0.0584, 0.0633, 0.0606, 0.0645, and the ratio of blueing slopes is 1.80, 1.80, 1.90, 1.83, an average of 1.83.

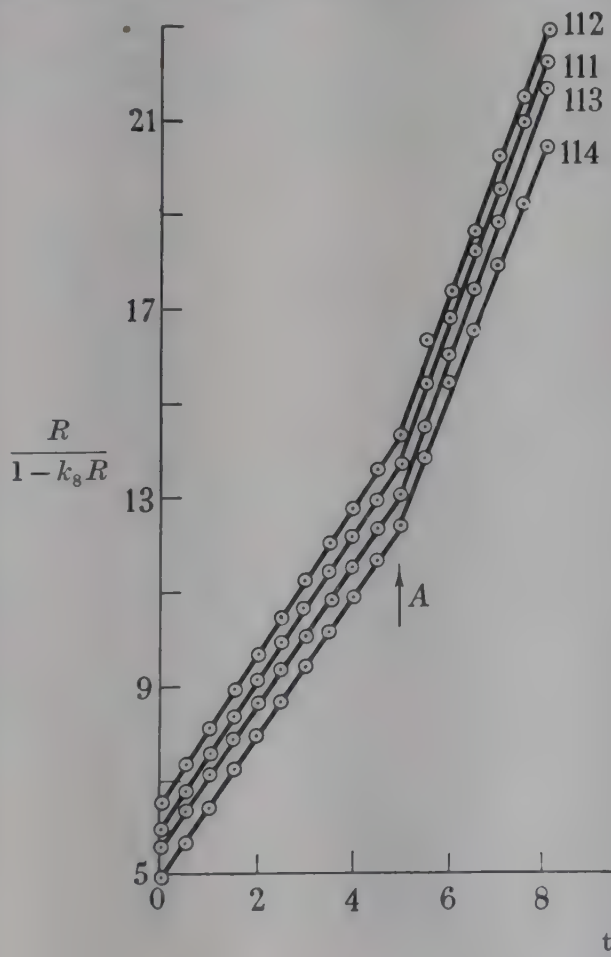


FIGURE 13. Butene-2 (*cis*): spacing 2.5 mm., 8.3 mm. H_2 : 0.41 mm. olefine at A.

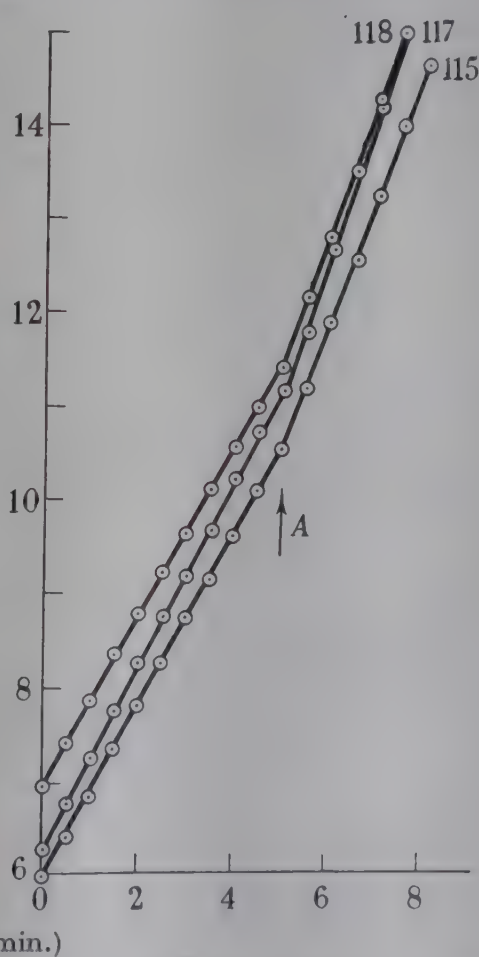


FIGURE 14. Butene-2 (*cis*): spacing 1.25 mm., 8.3 mm. H_2 : 0.41 mm. olefine at A.

Figure 14 shows runs 115, 117, 118 for a spacing of 1.25 mm. For these k_s is respectively 0.0650, 0.0688, 0.0649, and the blueing ratios are 1.46, 1.53, 1.51, an average value of 1.50.

From these values for the two spacings, b is 13.5 and the blueing coefficient 2.05. From model measurements, the diameter of the molecule is 6.5×10^{-8} cm.,

$$D = 8 \times 10^4 / (102.7 p_2 + 240 p_3)$$

$$= 84.3,$$

$$g = b^2 D = 1.54 \times 10^4,$$

$$k = g/n = 1.54 \times 10^4 / 0.41 \times 3.6 \times 10^{16}$$

$$= 1.04 \times 10^{-12} \text{ cm.}^3 \text{ molecules}^{-1} \text{ sec.}^{-1}.$$

$$\text{Collision number} = 1.5 \times 10^{-9}.$$

$$\text{Collision efficiency} = 1.04 \times 10^{-12} / 1.5 \times 10^{-9}$$

$$= 6.9 \times 10^{-4}.$$

Cyclo-hexene

Figure 15 shows runs 65, 67, 68 for molybdenum oxide at a spacing of 2.5 mm. The values of k_8 were 0.0625, 0.0665, 0.0751, the blueing ratios being 1.87, 1.93, 1.93, an average of 1.91.

Figure 16 shows runs 69, 70, 72 for a spacing of 1.25 mm., k_8 being 0.0826, 0.0750, 0.0879, the blueing ratios being 1.56, 1.52, 1.55, an average of 1.54.

From these, $b = 13$ and the blueing coefficient is 2.16.

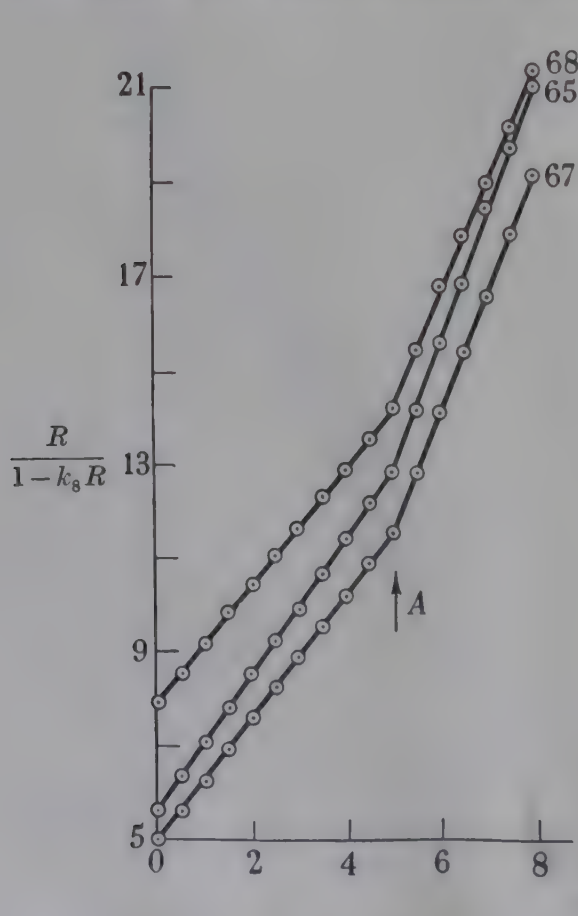


FIGURE 15. Cyclo-hexene: spacing 2.5 mm., 8.3 mm. H_2 : 0.41 mm. cyclo-hexene at A.

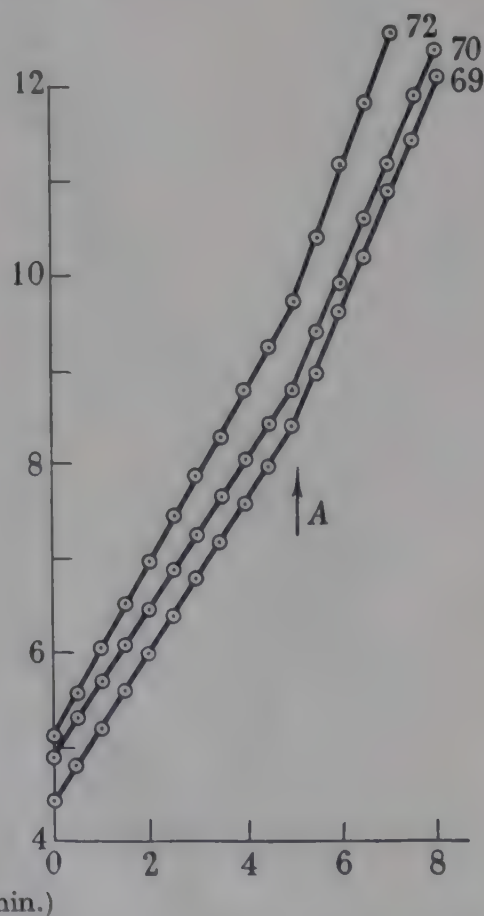


FIGURE 16. Cyclo-hexene: spacing 1.25 mm., 8.3 mm. H_2 : 0.41 mm. cyclo-hexene at A.

The estimated molecular diameter is 7×10^{-8} cm., giving

$$\begin{aligned} D &= 8 \times 10^4 / (102.7 p_2 + 266 p_3) \\ &= 83.4, \\ g &= b^2 D = 169 \times 83.4 = 1.41 \times 10^4, \\ k &= g/n \\ &= 9.6 \times 10^{-13}. \end{aligned}$$

The collision number is calculated to be 1.67×10^{-9} .

$$\begin{aligned} \text{Collision efficiency} &= 9.6 \times 10^{-13} / 1.67 \times 10^{-9} \\ &= 5.7 \times 10^{-4}. \end{aligned}$$

Benzene

Figure 17 shows runs 73, 74, 75 for molybdenum oxide at a spacing of 2.5 min. The values of k_8 are 0.0682, 0.0616, 0.0666 and the blueing ratios 1.22, 1.19, 1.24, an average of 1.22.

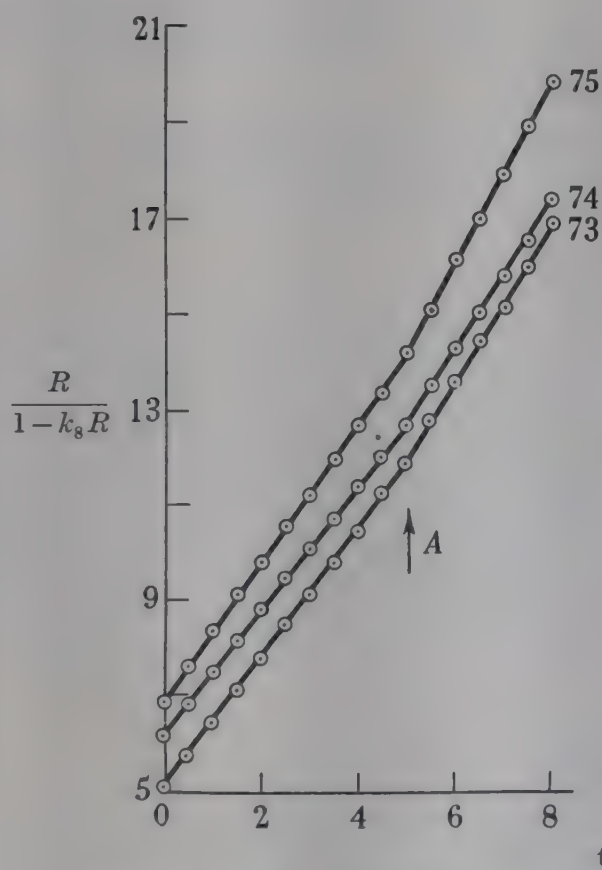


FIGURE 17. Benzene: spacing 2.5 mm., 8.3 mm. H_2 : 0.41 mm. benzene at A.

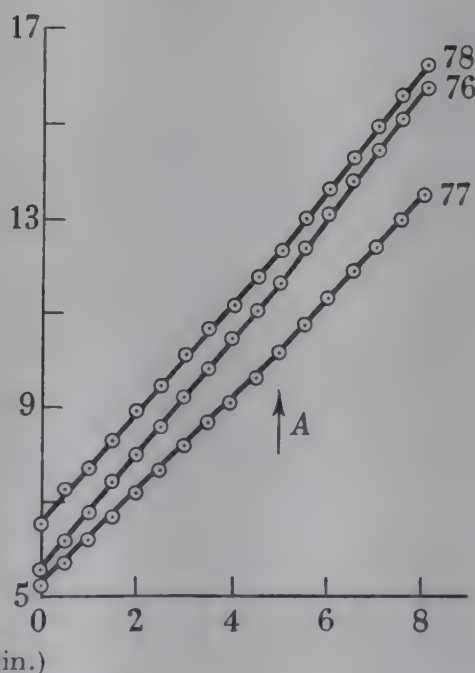


FIGURE 18. Benzene: spacing 1.25 mm., 8.3 mm. H_2 : 0.41 mm. benzene at A.

Figure 18 shows runs 76, 77, 78 for a spacing of 1.25 min., where k_8 is 0.0664, 0.0710, 0.0625 and the blueing ratios are 1.12, 1.14, 1.13, an average of 1.13.

From these, $b = 13$ and the blueing coefficient is 1.30.

The diameter of the benzene molecule is taken as 5.7×10^{-8} cm. and from this

$$D = 8 \times 10^4 / (102.7 p_2 + 203 p_3) \\ = 85.7,$$

$$g = b^2 D = 1.45 \times 10^4,$$

$$k = g/n = 1.45 \times 10^4 / 0.41 \times 3.6 \times 10^{16} \\ = 9.8 \times 10^{-13}.$$

The collision number is calculated to be 1.36×10^{-9} .

$$\text{Collision efficiency} = \frac{9.8 \times 10^{-13}}{1.36 \times 10^{-9}} = 7.2 \times 10^{-4}.$$

Cyclo-hexane

In view of the rather remarkable reactivity of benzene, it was considered advisable to try the effect of cyclo-hexane on the blueing rate, just to check up that all was well with these runs with cyclic compounds and that no peculiar effects could be obtained with saturated compounds.

Figure 19 shows run 79, with molybdenum oxide at a spacing of 2.5 mm. The value of k_8 was 0.093, and no distinguishable difference in the rate of blueing could be observed when the cyclo-hexane was introduced.

This result shows that no dehydrogenation reaction is occurring and that the results obtained above are truly functions of the unsaturated nature of the compounds.

Acetylene

Figure 20 shows two runs with acetylene. Run 127 was done with molybdenum oxide at a spacing of 2.5 mm.; the value of k_8 is 0.0681, and no detectable change in slope is observed on addition of the acetylene. To check that this effect is not the same as that observed with 2.3.3 trimethyl butene-1 where the radical blueing effect was the same as that for the alkyl radical produced, a check run, 128, was done with tungsten oxide at a spacing of 2.5 mm. The value of k_8 here was 0.0804 and again no detectable change in slope was observed, thus demonstrating that the efficiency of interaction between a hydrogen atom and acetylene is not nearly so high as that between a hydrogen atom and an olefine.

Butadiene

Figure 21 shows runs 119, 120, 121 for molybdenum oxide at a spacing of 2.5 mm. The values for k_8 are 0.0613, 0.0526, 0.0605. The blueing portion of the run after addition of butadiene shows a characteristic initial fast portion followed by a linear relationship. The ratio of the slope of this linear part to the pure hydrogen rate is 1.36, 1.40, 1.36.

Figure 22 shows runs 124, 125 for a spacing of 1.25 mm. The values of k_8 are 0.0680, 0.0685. Again, the rapid initial portion is present and the linear parts give ratios 1.65, 1.56.

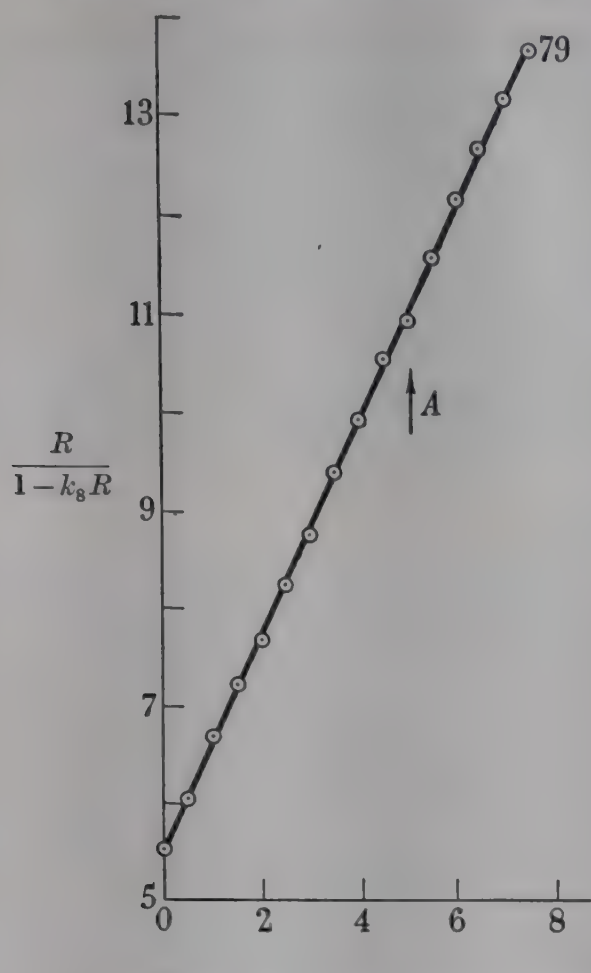


FIGURE 19. Cyclo-hexane: spacing 2.5 mm., 8.3 mm. H_2 : 0.41 mm. hexane at A .

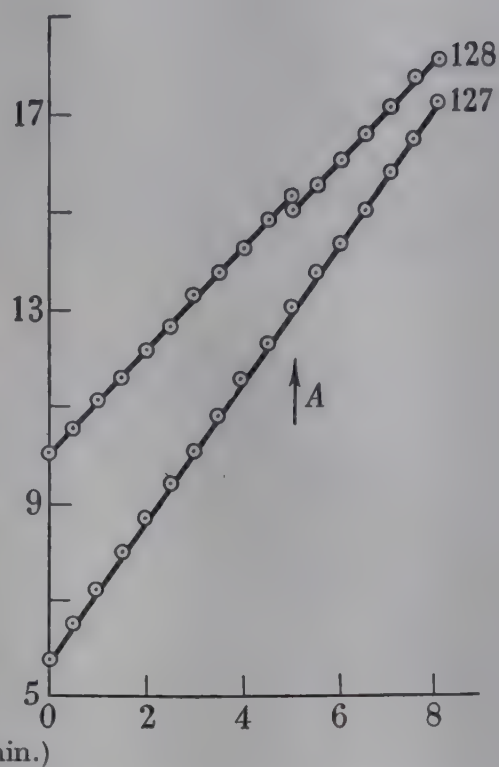


FIGURE 20. Acetylene: spacing 2.5 mm., 8.3 mm. H_2 : 0.41 mm. acetylene at A .

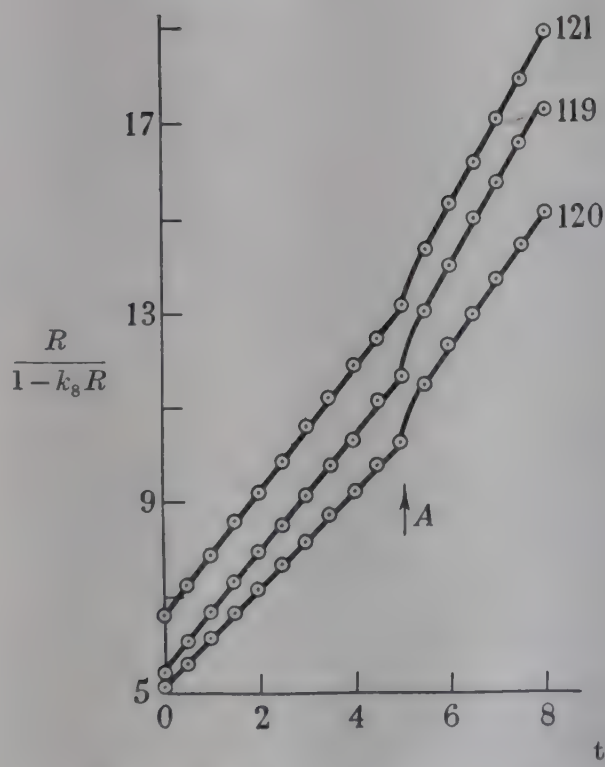


FIGURE 21. Butadiene: spacing 2.5 mm., 8.3 mm. H_2 : 0.41 mm. diene at A .

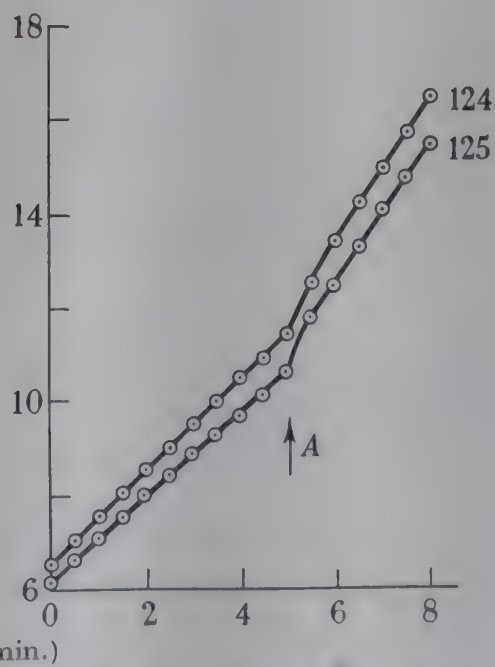


FIGURE 22. Butadiene: spacing 1.25 mm., 8.3 mm. H_2 : 0.41 mm. diene at A .

This curious result is explicable if one assumes that, at the greater separation distance, since the radicals will have a longer life-time before removal of the oxide, the possibility of the radical undergoing further reaction is increased. This may lead to the radical losing a hydrogen atom and returning to butadiene, or undergoing reaction with a hydrogen molecule, thus forming a stable butene-1 or butene-2 molecule and generating a fresh hydrogen atom. A further possibility is that the C_4H_7 radical may be removed on the oxide and subsequent hydrogen atoms reaching the oxide surface may, instead of adding to the oxide, be removed by addition to the remaining double bond of the C_4H_7 radical. These possibilities lead to a very complex situation which it is not possible at present to separate into well-defined simple steps.

Oxygen

A few qualitative experiments have been done on the interaction of atomic hydrogen with oxygen. Air, freed from condensable gases by passing through liquid air traps was used, the purpose of the experiment being to investigate whether or not the blue molybdenum oxide could be transformed to the white by this means. It was found that a slight removal of the blueing could be accomplished in presence of oxygen, but the process could not be taken very far. The striking feature of these experiments was that even small amounts of oxygen of the order of 0.07 mm. could completely stop the blueing process. Under these conditions, the blueing of the molybdenum oxide is completely stopped for a period of about 4 min. The rate then increases slowly to about the original blueing rate with hydrogen alone. This would seem to indicate a very efficient reaction between hydrogen atoms and an oxygen molecule. The only products that seem likely are HO_2 or OH and O . Although, from measurements made, the radicals formed would seem to have a positive, but small blueing coefficient, yet the observation that the blue colour can be removed would seem to indicate that a negative blueing coefficient is possible. This suggests perhaps a more complex picture in which the nature of the radicals formed depends on the surrounding conditions, such as plate separations, etc.

Further work must be done before the efficiency of $H + O_2$ interaction is obtained. Suffice it to say then, that the reaction is much more efficient than has hitherto been supposed. Under the conditions of experiment, it is highly improbable that a three-body collision could occur as has been usually postulated for $H + O_2$.

DISCUSSION

Table 1 contains the results obtained for the unsaturated compounds studied. For the sake of completeness those results previously obtained are also included. The most striking feature of these results is the remarkable similarity of the reactivities of these olefines, a reactivity which is not markedly affected by changing the structure in every possible way. Thus the main new principle that is established by this investigation is that the reactivity of the olefinic bond is practically independent of the structure. This generalization should be of importance in dealing with the reaction of radicals and olefines.

Recently, confirmation of this generalization has been obtained by the published results of two research groups working on very different experimental lines and under

different experimental conditions. The interaction of methyl radicals with olefines has been studied (Raal, Danby & Hinshelwood 1949) in the gaseous phase, the methyl radicals being obtained by the photo-decomposition of acetaldehyde. The temperatures employed are much higher than those in this investigation, being of the order 200 to 350° C. The kinetic scheme is somewhat complicated in that there is simultaneously occurring the chain decomposition of acetaldehyde, the chain polymerization of the olefine, termination by the mutual interaction of methyl radicals, by spontaneous isomerization of olefine polymer molecules and chain transfer by spontaneous decomposition of the growing polymer chain to form a methyl radical and a dead polymer molecule. The efficiency of interaction of a methyl radical with

TABLE 1

olefine	$10^{13} \times \text{velocity constant, } k,$ ($\text{cm.}^3 \text{ molecules}^{-1} \text{ sec.}^{-1}$)	$10^4 \times \text{collision}$ efficiency
ethylene	9.1	8.3
propylene	2.2	1.4
?-butene	7.0	4.7
<i>n</i> -butene-2 (<i>cis</i>)	5.8	2.8
2, 3, 3, tri-methyl-butene-1	9.9	5.2
tri-methyl-ethylene	9.5	5.0
	11.1*	5.8
tetra-methyl-ethylene	6.9	3.6
butene-2 (<i>trans</i>)	9.5	5
butene-2 (<i>cis</i>)	10.0	6.9
cyclo-hexene	9.6	5.7
benzene	9.8	7.2

* Determined by two different techniques.

the olefine double bond is compared with the efficiency of the interaction of a methyl radical with acetaldehyde. The general principle here shown is also that the reactivity of olefinic double bond is practically unchanged by substitution. Another interesting point in this paper is that ethylene is the least reactive olefine of the series ethylene, propylene and butene, having about half the reactivity of the others. The results presented in our paper indicate that propylene is less reactive than the others. In a previous paper (Danby & Hinshelwood 1941) it is shown that for methyl radicals, produced by the photolysis of acetaldehyde, propylene is more reactive than ethylene, whereas for ethyl radicals from the photolysis of propionaldehyde, ethylene is more reactive than propylene. This may be due to specific radical interference in the acetaldehyde decomposition or, alternatively, to the simple fact that although olefinic double bond reactivities are very similar, they may be modified by the nature of the attacking radical.

The other investigation referred to above (Kharasch & Sage 1949) has been conducted in the liquid phase. By the photolysis of tri-chloro-bromo-methane, these workers have produced free trichloro-methyl radicals. These are allowed to interact with olefinic double bonds, the interaction resulting in the addition of the $\text{C} \cdot \text{Cl}_3$ radical to the double bond. The radical so formed interacts with another $\text{C} \cdot \text{Cl}_3 \cdot \text{Br}$ molecule, abstracting the bromine atom so propagating the chain. The values for

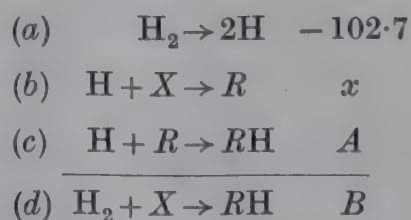
reactivities quoted are not absolute, but have been obtained relatively by having two double bonded compounds present, each competing for the $\text{C} \cdot \text{Cl}_3$ radical. A few of the results have been abstracted from this paper and are shown in table 2. Here again, we have further support for the generalization that the reactivity of the olefinic double bond is modified only slightly by changing the structure of the molecule.

TABLE 2. RELATIVE REACTIVITIES OF OLEFINIC DOUBLE BONDS
TOWARDS THE CCl_3 RADICAL

compound	reactivity
2-ethyl 1-butene	1.4
1-octene	1.0
2-methyl 2-butene	0.9
vinyl acetate	0.8
allyl-benzene	0.5
4, 4, 4, trichloro 1-butene	0.3
cyclo-hexene	0.2

For ethylene, from the results presented in Table 1, it can be shown that the maximum possible value for the energy of activation, assuming a steric factor of 1, is 4.25 kcal. From theoretical considerations, Evans & Szwarc (1949) assign a value of 0.025 for P in the reaction $X + \text{C}_2\text{H}_4$. This leads to an energy of activation of 2.04 kcal. For $X + \text{C}_3\text{H}_6$ their value is 0.01, leading to an activation energy of 2.55 kcal. It is unlikely that for the higher hydrocarbons P would change significantly from 10^{-2} . The average value for activation energy over the whole olefine series quoted in table 1, computed on the above basis is 2.1 ± 0.5 kcal.

It has been shown by various workers that in a series of related reactions, the energy of activation of the reaction is related to the heat of reaction in a simple manner. The difference in activation energy of two kinetically similar reaction steps is found to be of the order 0.2 to 0.25 times the change in heat of reaction (Evans & Polanyi 1936, 1938; Butler & Polanyi 1943; Steiner & Watson 1947). This proportionality constant is given the symbol β and the value used here will be 0.22. The heat of addition of a hydrogen atom to a double bond can be calculated by considering the following scheme.

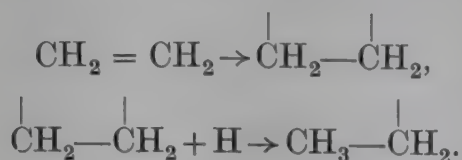


X is the olefine molecule and R the alkyl radical produced. The heat of reaction (c) is obtained by a knowledge of the $R\text{--H}$ bond strength and the values used here are those quoted by Steacie (1946). The values for the small heats of hydrogenation in reaction (d) are obtained from Conant & Kistiakowsky (1937). The calculations are listed in table 3.

TABLE 3. CALCULATION OF HEAT OF ADDITION OF A HYDROGEN ATOM TO AN OLEFINE

olefine	A (kcal.)	B (kcal.)	x
ethylene	+ 97.5	+ 32.6	37.8
propylene	+ 89.0	+ 29.9	43.6
<i>i</i> -butene	+ 86.0	+ 28.2	44.9
butene-2 (<i>cis</i>)	+ 89.0	+ 28.2	41.9
butene-2 (<i>trans</i>)	+ 89.0	+ 27.2	41.9
trimethyl ethylene	+ 86.0	+ 26.7	43.4
tetramethyl ethylene	+ 86.0	+ 26.4	43.1

A further method of obtaining the heat of reaction of $H + \text{olefine} \rightarrow \text{alkyl radical}$ is by considering the sequence, illustrated for ethylene below:



Values for opening the second valence bond in an olefine can be obtained from the calculations of Baughan & Polanyi (1940). The values used are given in table 4.

TABLE 4. ENERGY REQUIRED TO OPEN UP SECOND VALENCE OF DOUBLE BOND

	A (kcal.)
ethylene	56.6
propylene	53.0
<i>i</i> -butene	51.4
butene-2	49.4
trimethyl ethylene	47.8
tetramethyl ethylene	46.2

Further, from the same work, the value for the energy of a C—H bond can be obtained making the justifiable assumption that the bond energy is dependent only on the atoms which are attached to the carbon atom in question and not on the influence of groups more remotely situated. These bond energies are given in table 5. From the values in tables 4 and 5 the heat of addition of a hydrogen atom to any olefine can be calculated. The values obtained in this way can be seen in table 6. By taking $\beta = 0.22$, a value for the difference in activation energy between any of these olefines can be obtained since $\Delta E = \beta \Delta Q$. The steric factors calculated by Evans & Szwarc

(1949) can then be used to obtain the value of $\frac{P_2}{P_1} e^{\Delta E/RT}$, where P_2 is the steric factor for the olefine, P_1 for ethylene, the reference olefine and ΔE the difference in activation energy between the two. Since ethylene has two equivalent ends, it is necessary to multiply the calculated reactivity by 2. It can be shown that the least reactive end of the propylene double bond has a reactivity of 0.16 relative to the most reactive end. For *i*-butene, the factor is 0.4. For butene-2 and tetramethyl ethylene, the factor is again 2 and for trimethyl ethylene it is 0.25. Applying all these corrections for reactivity, the values in table 7 are obtained.

The columns in table 7 are practically self-explanatory. *R* is the calculated reactivity relative to ethylene after all the above-mentioned factors have been taken into consideration.

TABLE 5. EFFECT OF SUBSTITUENTS ON C—H BOND ENERGY

group	bond energy (kcal.)
$\begin{array}{c} \text{H} \\ \\ \text{H}-\text{C}\cdots\cdots\text{H} \\ \\ \text{H} \end{array}$	103.6
$\begin{array}{c} \text{H} \\ \\ -\text{C}-\text{C}\cdots\cdots\text{H} \\ \quad \\ \quad \text{H} \end{array}$	96.4
$\begin{array}{c} \quad \quad \\ \quad \quad -\text{C}- \\ \quad \\ -\text{C}-\text{C}\cdots\cdots\text{H} \\ \\ \text{H} \end{array}$	91.0
$\begin{array}{c} \quad \quad \\ \quad \quad -\text{C}- \\ \quad \\ -\text{C}-\text{C}\cdots\cdots\text{H} \\ \quad \\ \quad \quad -\text{C}- \end{array}$	86.9

TABLE 6. CALCULATED HEAT OF ADDITION OF A HYDROGEN ATOM TO AN OLEFINE

olefine	heat of addition (kcal.)
ethylene	39.8 (37.8)*
propylene	43.4 (43.6)
<i>i</i> -butene	45.0 (44.9)
butene-2	41.6 (41.9)
trimethyl ethylene	43.2 (43.4)
tetramethyl ethylene	40.7 (43.1)

* The values in parenthesis are those from table 3 for comparison.

TABLE 7. CALCULATED REACTIVITIES OF OLEFINES

olefine	thermochemical values						bond energy values					
	ΔQ	ΔE	<i>P</i>	$\frac{P_2}{P_1} e^{\Delta E/RT}$	<i>R</i> (calc.)	<i>R</i> (exp.)	ΔQ	ΔE	<i>P</i>	$\frac{P_2}{P_1} e^{\Delta E/RT}$	<i>R</i> (calc.)	<i>R</i> (exp.)
ethylene	0	0	0.025	1	1	1.0	0	0	0.025	1	1	1.0
propylene	5.8	1.28	0.01	3.36	1.95	0.24	3.6	0.79	0.010	1.48	0.86	0.24
<i>i</i> -butene	7.1	1.56	0.008	4.32	2.25	0.77	5.2	1.14	0.008	1.76	0.92	0.77
trimethyl ethylene	5.6	1.23	0.008	2.48	1.55	1.1	3.4	0.75	0.008	1.12	0.70	1.1
butene-2 (<i>cis</i>)	4.1	0.90	0.008	1.44	1.44	1.1	1.8	0.40	0.008	0.64	0.64	1.1
butene-2 (<i>trans</i>)	3.1	0.68	0.008	1.00	1.00	1.04						
tetramethyl ethylene	5.3	1.17	0.008	2.22	2.22	0.76	0.9	0.20	0.008	0.44	0.44	0.76

The values are so close together that the predicted order of reactivity as calculated could not be expected, not only because the experimental accuracy is not sufficiently great, but also because the bond energies are known to about ± 1 kcal. This sets the limiting accuracy of ΔQ at $\pm 25\%$ at best. Until more accurate values are known for bond strengths, it is impossible to do this type of calculation with any certainty.

However, it will be seen from table 7 that the overall change in reactivity is by a factor of about 2.5 and from the experimental values in table 1, the variation is by a factor of about 3.5. The close agreement between experiment and theory is considered to be very satisfactory, not only in supporting the experimental evidence, but in justification of this type of theoretical calculation.

From a calculation of the absolute value of the energy of activation for the reaction $\text{H} + \text{C}_2\text{H}_4 \rightarrow \text{C}_2\text{H}_5$ in the manner outlined by Evans, Gergely & Seaman (1948), a value of 5.3 kcal. is obtained, but since the computed value for the resonance energy of the transition state complex is 36 kcal., even relatively small changes in this exert a great influence on the value for the energy of activation. However, this result indicates the order of magnitude of the energy and is in reasonably good agreement with the possible value deduced from experiment.

With regard to the results for benzene, it seems rather surprising that this molecule should have a reactivity comparable to that of an olefine. However, it has been shown (Stein & Weiss 1948) that if radicals are formed from water by means of ionizing radiations, subsequent interaction between these radicals and benzene, suspended in the medium, is observed. Further qualitative supporting information has been obtained concerning this reaction (Forbes & Cline 1941) whereby the interaction of atomic hydrogen with benzene in the gaseous phase is reported as giving rise to reasonable quantities of hydrogenated diphenyls and cyclo-hexa-diene. This makes it a reasonable assumption that the primary step is addition of a hydrogen atom to the benzene molecule. If one takes into account the 6-fold symmetry for benzene, the specific rate constant per point of attack is 1.6×10^{-13} compared with 4.5×10^{-13} for ethylene. This is to be understood in terms of loss of resonance energy by the benzene molecule.

The lack of reactivity shown by acetylene under these conditions is somewhat surprising and it is not yet possible to state whether this is due to unreactivity of the acetylene, extreme instability of the radical produced, or a subsequent very fast reaction between the radical and a hydrogen molecule, resulting in the formation of a molecule of ethylene and regenerating a hydrogen atom. The most probable explanation is that the initial addition of the hydrogen atom to acetylene is not so efficient as to an olefine double bond.

The authors are indebted to Philips Petroleum Co., Texas, for samples of butene-2 *cis* and *trans*; to Anglo-Iranian Oil Co., for a sample of tri-methyl ethylene and to Bataafsche Petroleum Maatschappij for tetra-methyl ethylene.

One of us (J. C. R.) wishes to thank the Department of Scientific and Industrial Research for a Senior Research Award (1948-50) during the tenure of which these investigations were carried out.

REFERENCES

- Baughan, E. C. & Polanyi, M. 1940 *Nature*, **146**, 685.
 Butler, E. T. & Polanyi, M. 1943 *Trans. Faraday Soc.* **39**, 19.
 Conant, J. B. & Kistiakowsky, G. B. 1937 *Chem. Rev.* **20**, 181.
 Danby, C. J. & Hinshelwood, C. N. 1941 *Proc. Roy. Soc. A*, **179**, 169.
 Evans, M. G. & Polanyi, M. 1936 *Trans. Faraday Soc.* **32**, 1933.
 Evans, M. G. & Polanyi, M. 1938 *Trans. Faraday Soc.* **34**, 22.
 Evans, M. G., Gergely, J. & Seaman, E. C. 1948 *J. Pol. Sci.* **3**, 866.
 Evans, M. G. & Szwarc, M. 1949 *Trans. Faraday Soc.* **45**, 940.
 Forbes, G. S. & Cline, J. E. 1941 *J. Amer. Chem. Soc.* **63**, 1713.
 Kharasch, M. S. & Sage, M. 1949 *J. Organ. Chem.* **14**, 537.
 Melville, H. W. & Robb, J. C. 1949 *Proc. Roy. Soc. A*, **196**, 445, 466, 479, 494.
 Raal, F. A., Danby, C. J. & Hinshelwood, C. N. 1949 *J. Chem. Soc.* no. 475, p. 2219.
 Steacie, E. W. R. 1946 *Atomic and free radical reactions*. New York: Reinhold.
 Stein, G. & Weiss, J. 1948 *Nature*, **161**, 650.
 Steiner, H. & Watson, A. R. 1947 *Faraday Soc. Disc.* **2**, 88.

Statistical mechanics of the two-dimensional assembly

BY H. N. V. TEMPERLEY

King's College, University of Cambridge

(Communicated by D. R. Hartree, F.R.S.—Received 12 January 1950—
 Revised 15 March 1950)

The work of Kaufman & Onsager (1946) on the two-dimensional Ising model of a ferromagnet is extended from the plane square lattice to the plane honeycomb and triangular lattices. The specific heat anomaly, where it exists, turns out to be of the same type in all three lattices, an infinity in the specific heat at the Curie temperature. It is concluded that second-nearest neighbour interactions may have a considerable effect on the position of the Curie temperature.

1. INTRODUCTION

The problem of evaluating the exact partition function for the two-dimensional square lattice was solved by Onsager (1944) for the case of a network of atomic magnets at fixed distances apart, each magnet having two possible orientations only, all interactions except those between nearest neighbours being ignored. A very much shorter method of treating this problem was found by Kaufman & Onsager (1946), and, in this paper, we shall apply the latter method to the two-dimensional honeycomb and triangular lattices. It is only necessary to make the calculation for the honeycomb lattice, as the result for the triangular lattice can at once be inferred from it, by methods described by Wannier (1945).

2. STATEMENT OF THE PROBLEM IN TERMS OF THE THEORY OF OPERATORS

For the plane square lattice it is known (see, for example, Onsager (1944)) that the problem is equivalent to the evaluation of the eigenvalues of the following operator:

$$V = V_1 V_2 = \exp(H^* \sum_r C_r) \exp(H' \sum_r S_r S_{r+1}) \quad (1)$$

where H, H' , are the quantities $J/kT, J'/kT, 2J$ being the energy difference between a like pair of vertical nearest neighbours and an unlike pair, while $2J'$ is the same quantity for the horizontal pairs. The symbol * following any quantity means that it is to be modified by the following algebraic transformation:

$$H^* = \frac{1}{2} \log \coth H \text{ (the dual transformation).}$$

The quantities C_r, S_r are quaternion operators defined as follows: We first define a series of commuting square roots of unity μ_r . S_r is then the operation of multiplying the quantity that follows it by μ_r , while C_r is the operation of replacing μ_r by $-\mu_r$ wherever it appears, but leaving all the other μ 's unaltered, all algebraic numbers also being unaltered.

Kaufman & Onsager (1946) solve this problem by building up a set of anti-commuting quantities according to the following scheme:

$$\left. \begin{aligned} P_1 &= S_1, & Q_1 &= iC_1S_1, \\ P_2 &= C_1S_2, & Q_2 &= iC_1C_2S_2, \\ P_3 &= C_1C_2S_3, & Q_3 &= iC_1C_2C_3S_3, \\ &\dots\dots\dots, & &\dots\dots\dots \end{aligned} \right\} \quad (2)$$

Then the following relations hold

$$\left. \begin{aligned} V_1P_jV_1^{-1} &= \cosh 2H^*P_j - i \sinh 2H^*Q_j \\ V_1Q_jV_1^{-1} &= i \sinh 2H^*P_j + \cosh 2H^*Q_j, \\ V_2P_{j+1}V_2^{-1} &= \cosh 2H'P_{j+1} + i \sinh 2H'Q_j, \\ V_2Q_jV_2^{-1} &= -i \sinh 2H'P_{j+1} + \cosh 2H'Q_j, \end{aligned} \right\} \quad (3)$$

a transformation involving the operator C_r affecting only the quantities P_r, Q_r , while one involving the operator S_r, S_{r+1} affects only the quantities P_{r+1}, Q_r . A complication occurs because the quantity P_{r+n} , as defined above, n being the number of atoms in a horizontal row, is not equal to P_r . This causes a rather troublesome end-effect, but it can be shown rigorously (Kaufman 1949) that it can be neglected when n is very large. Kaufman also shows that, by virtue of the isomorphism between the spinor group and the rotation group, we can find the eigenvalues of operators such as (1) if we can solve the corresponding rotation problem. For practical purposes of calculating the partition function we need only the largest eigenvalue, as it has been shown independently by Kaufman (1949) and Temperley (unpublished), that for a very large assembly, no physical result is affected by neglecting the other eigenvalues.

We follow for the square lattice a method almost equivalent to that of Kaufman & Onsager (1946). Define new operators as follows

$$R_i = P_1 + \epsilon_i P_2 + \epsilon_i^2 P_3 + \dots, \quad T_i = Q_1 + \epsilon_i Q_2 + \epsilon_i^2 Q_3 + \dots, \quad (4)$$

where ϵ_i is the n th root of unity given by $\epsilon_i = \exp \left[\frac{2t\pi i}{n} \right]$. Then, neglecting the end-effect discussed above, transformation with the operator V_1 is equivalent to an imaginary rotation of the two operators R_i, T_i , while transformation with V_2 is

equivalent to an imaginary rotation of the two operators $R_i, \epsilon_i T_i$; the effect of the transformation of these two quantities by the whole operator (1) being described by a 2×2 matrix, whose eigenvalues are $e^{\pm \gamma_i}$, where

$$\cosh \gamma_i = \cosh 2H^* \cosh 2H' - \sinh 2H^* \sinh 2H' \cos \frac{2t\pi}{n}. \quad (5)$$

The largest possible eigenvalue of operator (1) is thus determined by multiplying together the quantities e^{γ_i} , t running from 1 to n , so that the logarithm of the largest eigenvalue can be expressed as a sum, which in the limit passes over into the expression

$$\int_0^{2\pi} \cosh^{-1} (\cosh 2H^* \cosh 2H' - \sinh 2H^* \sinh 2H' \cos \omega) d\omega, \quad (6)$$

which is the same as the integral found by Onsager (1944).

The singularity found by Onsager arises from the second derivative with respect to T of the integrand becoming infinite for ω zero, which results in the second derivative of the partition function becoming infinite, and this is only possible if H' and H are connected by the relation

$$\sinh 2H \sinh 2H' = 1, \quad (7)$$

which, for given values of the interactions J', J , can be satisfied by at most one value of T . It can be satisfied if the interactions are of the same sign (both favouring like pairs of nearest neighbours, or both favouring unlike pairs), but not if they are of opposite signs.

The plane honeycomb lattice may, without altering its topology, be deformed into a plane square lattice with one-quarter of the interactions omitted; thus we might include only the horizontal interactions between atoms 1 and 2, 3 and 4, 5 and 6, ... in the first row, the interactions between 2 and 3, 4 and 5, 6 and 7, ... in the second row, and so on alternately. It is possible to handle the case where all three interactions differ, by modifying the vertical interactions in the 'brick-wall' lattice into which our hexagonal lattice can be deformed. The problem can be reduced to that of finding the largest eigenvalue of the operator $V_1 V_2 V_3 V_4$, where

$$V_1 = \exp [H^*(C_1 + C_3 + C_5 + \dots) + H''*(C_2 + C_4 + C_6 + \dots)],$$

$$V_2 = \exp [H''(S_1 S_2 + S_3 S_4 + \dots)],$$

$$V_3 = \exp [H''*(C_1 + C_3 + C_5 + \dots) + H^*(C_2 + C_4 + C_6 + \dots)],$$

$$V_4 = \exp [H''(S_2 S_3 + S_4 S_5 + \dots)],$$

and H, H', H'' are the quantities J, J', J'' divided by kT , corresponding to the interactions in the three possible directions. The operation V_1 represents, apart from a factor, the effect of vertical interactions between atoms in the k th and $(k+1)$ th rows, of strength H between atoms bearing odd numbers and of strength H' between atoms bearing even numbers. (The atoms are numbered horizontally in rows, and corresponding numbers are vertically below one another.) The determination of the operators corresponding to vertical interactions is given by Onsager (1944). The operation V_2 represents the effect of horizontal interactions of strength H'' between atoms numbered 1 and 2, 3 and 4, etc., in the $(k+1)$ th row. The operation V_3 represents

the effect of vertical interactions between the $(k+1)$ th and $(k+2)$ th rows, of strength H' between atoms bearing odd numbers and of strength H between atoms bearing even numbers. The operation V_4 represents the effect of horizontal interactions between atoms numbered 2 and 3, 4 and 5, etc., in the $(k+2)$ th row. The vertical interactions between the $(k+2)$ th and $(k+3)$ th rows are again represented by V_1 , and so on all the way down the lattice, the effect of all the interactions being therefore represented by a high power of the operator $V_1 V_2 V_3 V_4$. If the lattice so formed is deformed into honeycomb shape, it will be found that the interactions H, H', H'' each correspond to one of the three possible directions in this lattice.

To determine the eigenvalues of this operator, we introduce the following quantities (similar to R_i and T_i as used in the treatment of the square lattice):

$$\left. \begin{aligned} U_i &= P_1 + \epsilon_i^2 P_3 + \epsilon_i^4 P_5 + \dots, \\ W_i &= P_2 + \epsilon_i^2 P_4 + \epsilon_i^4 P_6 + \dots, \\ X_i &= Q_1 + \epsilon_i^2 Q_3 + \epsilon_i^4 Q_5 + \dots, \\ Y_i &= Q_2 + \epsilon_i^2 Q_4 + \epsilon_i^4 Q_6 + \dots, \end{aligned} \right\}$$

where V_1, V_3 have the effect of rotating U_i and X_i , also W_i and Y_i , while V_2 rotates W_i and X_i and V_4 rotates U_i and $\epsilon_i^2 Y_i$. Thus, the problem reduces to that of evaluating the eigenvalues of a number of 4×4 matrices, which can be reduced to that of solving the following set of equations:

$$F^{(l)}(\lambda) = \lambda^4 - \left(\epsilon_l s^* + \frac{1}{\epsilon_l} s'^* \right) \lambda^3 - 2(s^* s'^* + c^* c'^* c'') \lambda^2 - \left(\frac{s^*}{\epsilon_l} + \epsilon_l s'^* \right) \lambda + 1 = 0. \quad (8)$$

where $c^* = \cosh 2H^*, \quad s'' = \sinh 2H'' \quad \text{etc.}$

Two roots of each of these equations are greater than unity numerically; let them be $\lambda_1^{(l)}, \lambda_2^{(l)}$, while the other two are numerically less than unity and have no effect on the partition function when raised to a high power, i.e. for a very large assembly. The logarithm of the largest eigenvalue of the operator $V_1 V_2 V_3 V_4$ is therefore proportional to the integral $\int_0^{2\pi} \log(\lambda_1(\omega) \lambda_2(\omega)) d\omega$ to which the series $\sum_l \log(\lambda_1^{(l)} \lambda_2^{(l)})$ reduces in the limiting case. Now we have identically

$$\lambda^{-2} F^{(l)}(\lambda) \equiv \lambda_1 \lambda_2 \left(1 - \frac{\lambda}{\lambda_1} \right) \left(1 - \frac{\lambda}{\lambda_2} \right) \left(1 - \frac{\lambda_3}{\lambda} \right) \left(1 - \frac{\lambda_4}{\lambda} \right).$$

Replacing λ by $-e^{i\omega'}$, taking logarithms of both sides, expanding the logarithms in convergent power series in $e^{i\omega'}$ or $e^{-i\omega'}$, it is evident that $\log(\lambda_1^{(l)} \lambda_2^{(l)})$ is the constant term in the Fourier expansion of $\log[e^{-2i\omega'} F^{(l)}(-e^{i\omega'})]$ as a function of ω' , that is

$$\frac{1}{2\pi} \int_0^{2\pi} \log[e^{-2i\omega'} F^{(l)}(-e^{i\omega'})] d\omega' = \log(\lambda_1^{(l)} \lambda_2^{(l)}),$$

and therefore

$$\int_0^{2\pi} \log(\lambda_1(\omega) \lambda_2(\omega)) d\omega = \frac{1}{2\pi} \int_0^{2\pi} \int_0^{2\pi} \log[e^{-2i\omega'} F^{(l)}(-e^{i\omega'})] d\omega d\omega'. \quad (9)$$

The part of $\log f$ involving an integral can now be reduced to the form, symmetrical between the three interactions,

$$\log f \propto \frac{1}{2\pi} \int_0^{2\pi} \int_0^{2\pi} \log [c^* c'^* c''^* + s^* s'^* s''^* - s^* \cos(\omega' + \omega) - s'^* \cos(\omega' - \omega) - s''^* \cos 2\omega'] d\omega d\omega', \quad (10)$$

and the result for the plane triangular lattice can be obtained by taking the dual of this result (see Wannier 1945). The result for the plane honeycomb lattice may be written in terms of the variables H, H', H'' themselves, the log (partition function) containing the integral

$$\int_0^{2\pi} \int_0^{2\pi} \log [cc'c'' + 1 - s's'' \cos(\omega' + \omega) - s's'' \cos(\omega' - \omega) - s's'' \cos 2\omega'] d\omega d\omega' \quad (11)$$

(I am much indebted to Professor L. Onsager for the derivation of equations (9) and (10) from equation (8).)

If all three interactions in the honeycomb lattice are equal, the calculations are much simplified, as equation (8) can be split into two quadratic equations whose roots can be written down explicitly, the largest root corresponding to a given value of t can at once be identified, and the logarithm of the largest eigenvalue can be calculated directly, as for the square lattice, the rather indirect method of deriving equation (10) not being required when the solution of equation (8) is possible in compact form. In the honeycomb lattice $\log f$ contains the integral

$$\int_0^{2\pi} \cosh^{-1} \left[\frac{1}{2} \left\{ \left(\frac{2c^3 + 2c^2}{s^2} + \cos^2 \omega \right)^{\frac{1}{2}} - \cos \omega \right\} \right] d\omega. \quad (12)$$

The result for the triangular lattice is best obtained from this by the dual transformation. For this lattice $\log f$ contains the integral

$$\int_0^{2\pi} \cosh^{-1} \left[\frac{1}{2} \left\{ \left(\frac{2c^3 + 2c^2 s}{s} + \cos^2 \omega \right)^{\frac{1}{2}} - \cos \omega \right\} \right] d\omega. \quad (13)$$

The values of the interaction energy for which the second derivative of the log grand with respect to temperature can become infinite are determined by the following equations: $\sinh 2H = \pm \sqrt{3}$ (honeycomb) and $\sinh 2H = 1/\sqrt{3}$ (triangular), compared with the equation $\sinh 2H = \pm 1$ (square). It is particularly to be noticed that, for the triangular lattice, there is no possible negative value of H , that is there is no Curie temperature when all three interactions are equal and favour unlike pairs of neighbours. This result is not surprising, as, in the triangular lattice, nearest neighbours of the same atom can be nearest neighbours of one another, which is not possible in the other two lattices.

Inspection of equations (12) and (13) shows that the specific heat anomaly, if it exists, is always a logarithmic infinity of Onsager type.

3. THE EFFECT OF SECOND NEAREST NEIGHBOUR INTERACTIONS IN THE PLANE SQUARE LATTICE

The triangular lattice is easily seen to be topologically equivalent to the square lattice with an interaction added along half of the diagonals, e.g. those sloping upwards from left to right. (The *true* second nearest neighbour model, with interactions along the other diagonals as well, seems to be a much more difficult problem. For example, it is not possible to find the dual of such a lattice, because some of the 'rods' cross one another.) In the square lattice, we suppose for simplicity that the horizontal and vertical interactions are both represented by H , and we can examine the effect of introducing an interaction along the diagonals represented by H'' by using the dual of expression (10), after putting H' equal to H , which reduces to

$$\frac{1}{2\pi} \int_0^{2\pi} \int_0^{2\pi} \log (C^2 C'' + S^2 S'' - 2S \cos \omega \cos \omega' - S'' \cos 2\omega') d\omega d\omega'. \quad (14)$$

It is easily shown that the condition for the integrand to become infinite is $\sinh 2H = e^{-2H''}$ for H positive, and $\sinh (-2H) = e^{-2H''}$ for H negative. Both these equations give a Curie temperature for all positive J'' , and for all negative J'' when J'' is numerically smaller than J .

The conclusion is therefore that second nearest neighbour interactions may have a very considerable effect on the location of a Curie point, and that, in the honeycomb and triangular lattices, the specific heat anomaly, where it occurs, is of the same type as occurs in the square lattice, namely, an infinity in the specific heat at the Curie temperature.

I wish to thank Professor L. Onsager and Dr R. Sack for very helpful discussions and correspondence, and I also wish to thank the Royal Society for the award of the Smithson Research Fellowship, during the tenure of which this work was carried out.

REFERENCES

- Kaufman, B. 1949 *Phys. Rev.* **76**, 1232.
 Kaufman, B. & Onsager, L. 1946 *Report of Cambridge Conference on fundamental particles and low temperatures*, **2**, 137.
 Onsager, L. 1944 *Phys. Rev.* **65**, 117.
 Wannier, G. H. 1945 *Rev. Mod. Phys.* **17**, 50.

Travelling disturbances in the ionosphere

By G. H. MUNRO

*Radio Research Board, Australian Commonwealth Scientific
and Industrial Research Organization*

*(Communicated by Sir Edward Appleton, F.R.S.—Received 13 January 1950—
Read 8 June 1950)*

Methods have been developed for the examination of the horizontal and vertical movements of short-period disturbances in the ionosphere.

It has been found that quasi-periodic travelling disturbances with periods of from 10 to 60 min. are of frequent occurrence in the *F* region by day. They appear as temporary variations in the vertical distribution of ionization which show a horizontal progression and a vertical progression downwards.

The horizontal directions of travel have a well-defined mean direction on most days. The mean direction shows a marked seasonal variation with a sudden change at each equinox.

The horizontal rate of travel is usually between 5 and 10 km./min., and the rate of vertical progression downwards is approximately half the horizontal rate.

The disturbances are considered to be variations of a compressional type in the atmosphere resulting in changes in the distribution of ionization.

1. INTRODUCTION

As part of the investigations of the Radio Research Board at Sydney in the years 1937–9 continuous recording of the virtual height and intensity of reflexions from the ionosphere was being carried out in the daytime on two frequencies close to the vertical penetration frequency of the *F* region. It was noticed that the penetration frequency often underwent fluctuations of a quasi-periodic nature with periods of the order of 10 to 30 min.

A variable-frequency ionosphere recorder was then in operation at Mount Stromlo Solar Observatory, Canberra, and the records from it were examined by Mr A. J. Higgs. He found that similar variations in critical frequency occurred at Canberra, and that they were followed by changes in virtual height which appeared first at the high frequencies and progressed to the lower frequencies.

When the Canberra records were compared with those at Sydney it was found that the changes usually occurred earlier at Canberra than at Sydney, the time difference being of the order of 15 to 30 min. At this stage the investigation was interrupted by the diversion of staff to war work.

In 1947, when fundamental research in this field was resumed, a system of spaced observing points was established near Sydney with the object of detecting horizontal movements in the ionosphere. It was again found that short-period disturbances in the *F* region were of frequent occurrence and that a disturbance appeared always to show a progression across the system of the spaced points, which was sufficiently systematic to enable deduction of a direction of movement and a rate of progression in a horizontal direction. A preliminary report on these movements was made in a letter to *Nature* (Munro 1948).

Similar deductions made from observations of a different type in England were reported at the same time by Beynon (1948). The only previous evidence of such movements appears to have been a deduction by Pierce & Mimno (1940) from observations of multiple echoes at night.

It was noticed that the periods of the variations observed with the spaced system on 5.8 Mc./sec. (mostly virtual height changes) were very similar to those observed earlier on higher frequencies.

A variable-frequency ionosphere recorder was, therefore, also used on a number of days to delineate the change of virtual height with frequency ($h'f$) at frequent intervals at a central point in the fixed-frequency observing network. Comparison of records then showed very clearly the sequence of events. A disturbance would become evident first as a fall in critical frequency and then as an increase of virtual height over part of the frequency range causing a 'kink' in the $h'f$ curve. This kink would move down the curve from the higher toward the lower frequencies on successive records. When it reached 5.8 Mc./sec. corresponding changes in virtual height, with the characteristic time differences, would be recorded at the spaced stations. An analysis of one such disturbance has been described (Munro 1949).

With this knowledge of the nature of the disturbances it was possible to compare the Sydney observations more closely with routine records from $h'f$ recorders at more distant points and so obtain further information on the extent and spatial variations of the disturbances.

The spaced system has now been in effectively continuous operation for over a year.

This paper is concerned mainly with the examination of experimental results obtained and their interpretation in terms of ionization changes.

2. OBSERVING SYSTEM AND EQUIPMENT

The stations in the observing system were located as indicated in figure 1. No. 1 is the base station at the Electrical Engineering Department of the University of Sydney (lat. $33^{\circ} 53' S$, long. $151^{\circ} 11' E$). No. 2 is at Cronulla, no. 3 at Camden, and no. 4 at Liverpool. No. 5, at Blaxland, was added more recently to replace Cronulla and give improved accuracy. The distances of these field stations from the University are 21, 50, 25 and 56 km. respectively.

Fixed-frequency pulse transmitters all operating on the same frequency of 5.8 Mc./sec. were located at points 1, 2, 3, and later at point 5. They are of conventional type, giving a peak power of 1 kW and a pulse duration of $30 \mu\text{sec.}$ at a repetition frequency of 50 c./sec. The transmitters at points 2 and 3 operated automatically and were triggered through a controllable delay circuit by the pulses from the transmitter at point 1 received by the direct path. This ensured positive synchronization of repetition frequency.

Recording has been carried out mainly at point 1, with supplementary records at point 3. At each receiving point the radiation from the local transmitter will be returned vertically (assuming a plane horizontal ionosphere), but that from each of the two remote transmitters will be reflected at a point directly above the mid-

point between the transmitters and the receiver and will, thus, be received at an angle a few degrees off the vertical. Comparison of records at point 1, where the reflected ray from T_1 is vertical and that from T_3 oblique, and at T_3 , where the reverse is the case, has shown that the effects of this slight obliquity are negligible, and also that any effects due to lateral deviation in the ionosphere are the same at all three points within the accuracy of observation.

The triangles across which movement is observed are therefore regarded as those shown by the heavy lines in figure 1. The system of three spaced transmitters and a recorder at one of the transmitting points will be referred to as a 'three-point system'.

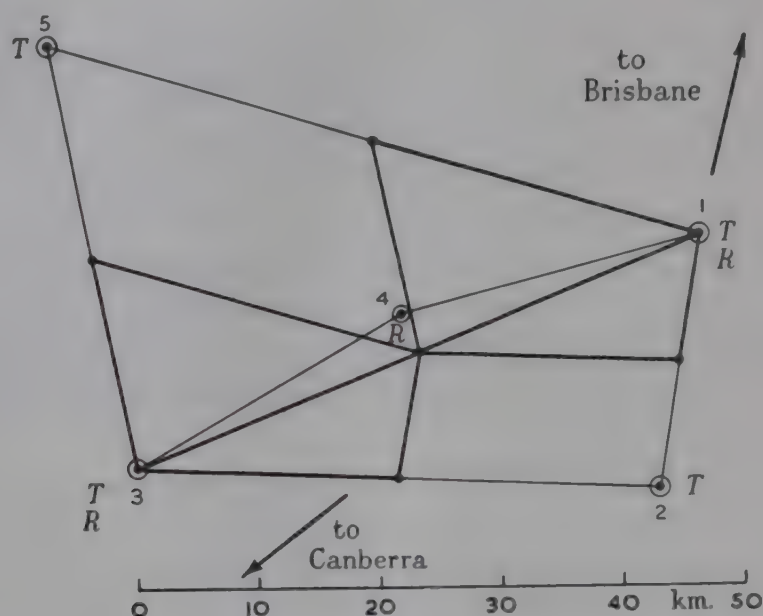


FIGURE 1. Plan of observing system.

The technique follows customary modern procedure. Receivers and cathode-ray oscillograph display units have been designed to reproduce the short pulses without serious distortion. An adjustable time-base sweep and height calibration marks are provided. Intensity modulation of the cathode-ray beam is used to present both height marks and the echo pattern on the screen, and the picture is photographed on continuously moving 35 mm. film. Periodical interruptions of the transmitters provide both time marks and transmitter identification on the record. The adjustable delay in the triggering system at the remote stations enables the pulses transmitted from the three stations to be separated by short intervals, so that all three echo patterns appear on the cathode ray tube trace allowing them to be recorded on the one film. This greatly simplifies detection of time differences of occurrence of events on the three sets of echoes. The pulses from the three stations are normally spaced about 300 μ sec. apart, corresponding to 50 km. on the height scale. A sample record is shown in figure 6a.

The variable-frequency ionosphere recorder was one of a standard type constructed by the Radiophysics Laboratory. It has a frequency range from 1 to 15 Mc./sec. swept in 2.5 min. It was located at point 4, and was operated continuously during the period of the tests made to provide comparisons between these

variable-frequency records and the fixed-frequency records on the three-point system.

Similar recorders are used for routine recording at Canberra and Brisbane, but they make a sweep only once every 10 min. A simple system which enabled recording of phase-path changes as well as virtual heights (group path) was in use most of the time at point 1, and for a short period at point 3 also.

3. CHARACTERISTICS OF A TYPICAL DISTURBANCE

The most obvious feature of a disturbance as observed at a fixed radio-frequency is a change in the virtual height of reflexion. On 5 July 1948 a well-defined example occurred, which appeared successively on the records at Canberra, Sydney and Brisbane in much the same form. The ordinary ray virtual heights for a frequency of 5.8 Mc./sec. at these three locations are plotted against time in figure 2. The

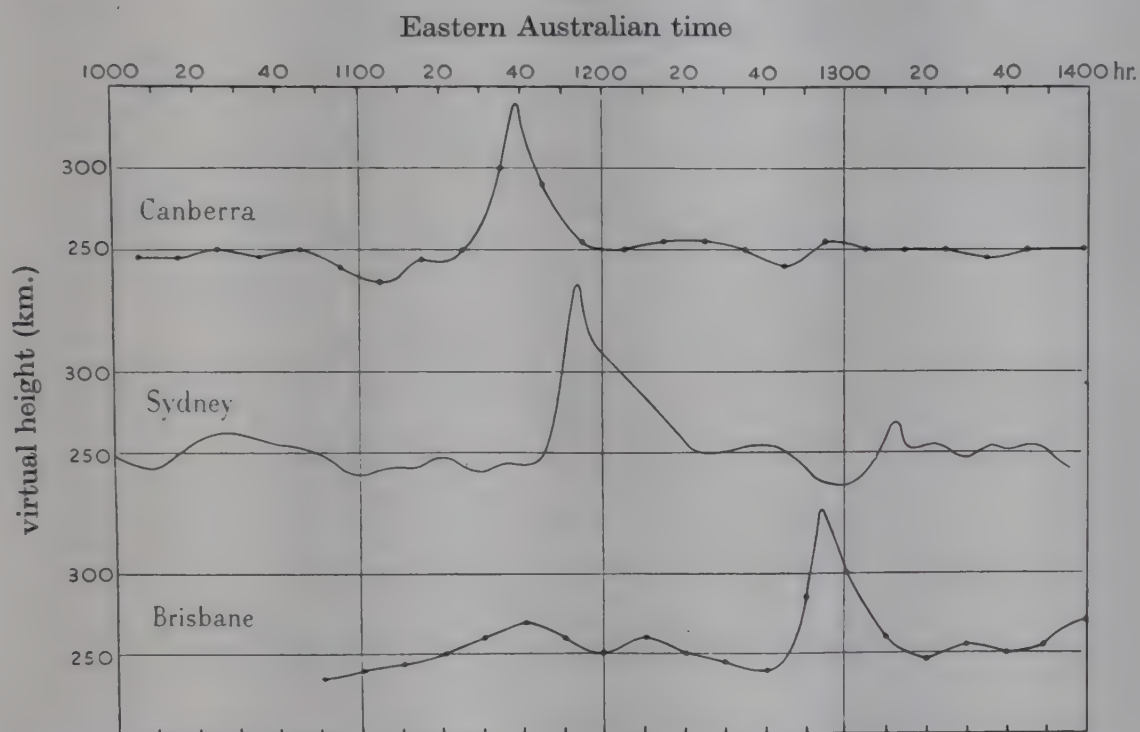


FIGURE 2. Virtual height variations of ordinary ray on 5.8 Mc./sec. 5 July 1948.

Canberra and Brisbane figures are from $h'f$ records taken at 10 min. intervals, so the height variations are somewhat smoothed, while the Sydney curve is copied from a continuous record and shows all the details of the changes. The direction of travel and apparent horizontal velocity of this disturbance were observed on the three-point system as being 2° east of north and 10 km./min. respectively. The apparent velocity deduced from the Canberra-Brisbane time difference when corrected for the direction of travel observed at Sydney agrees very closely with that obtained at Sydney.

It is seen, therefore, that such a disturbance can travel horizontally at least 900 km. without marked change in form or velocity.

The manner in which the disturbance affects the distribution of ionization in the F region may be seen from an examination of the $h'f$ curves. In figure 3 is shown a

set of tracings of the actual curves for both the ordinary and extraordinary rays as recorded at Brisbane during the passage of the disturbance. They show virtual height on a linear scale vertically, plotted against frequency on a logarithmic scale horizontally. The height scales are overlapped for compactness. Each time indicated refers to the curve immediately above the corresponding 200 km. height line. A vertical line indicates a frequency of 5.8 Mc./sec., so that changes in height occurring at that frequency may be examined.

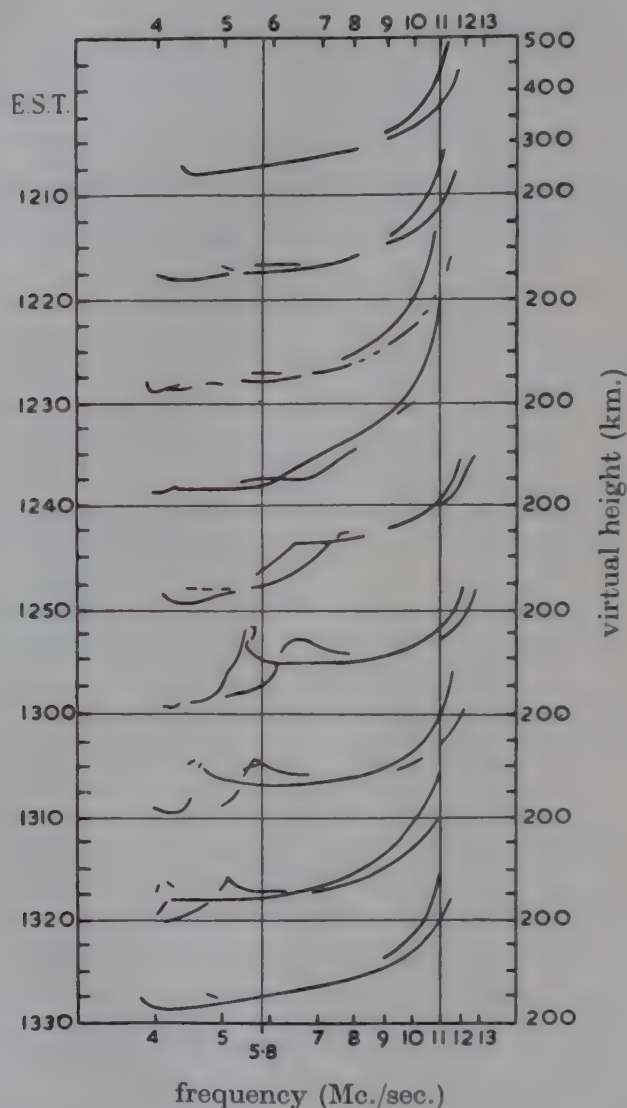


FIGURE 3. Tracings of Brisbane $h'f$ records. 5 July 1948.

The sequence of events is typical of such a disturbance. First, at 1210 there is a normal $h'f$ curve. Then the critical frequency of the ordinary ray falls slightly (1220), and the virtual height at frequencies near penetration increases markedly (1230). The increase in virtual height then shows lower down the curve (1240) and later becomes a distinct 'kink' in the curve. Meanwhile, the critical frequency has risen again (1250). It will now be observed also that similar changes are occurring in both rays but always later in the extraordinary ray, resulting in the characteristic loop and cross-over formation. The kinks appear successively farther down the curves until, at 1330, the curve is very similar to the first one at 1210.

The heights shown in figure 3 are virtual heights and, owing to group retardation, they will be considerably greater than the true heights of reflexion. It is, however, possible to derive the true height curves from the virtual height curves for the ordinary ray. This has been done, using the method outlined by Manning (1947).

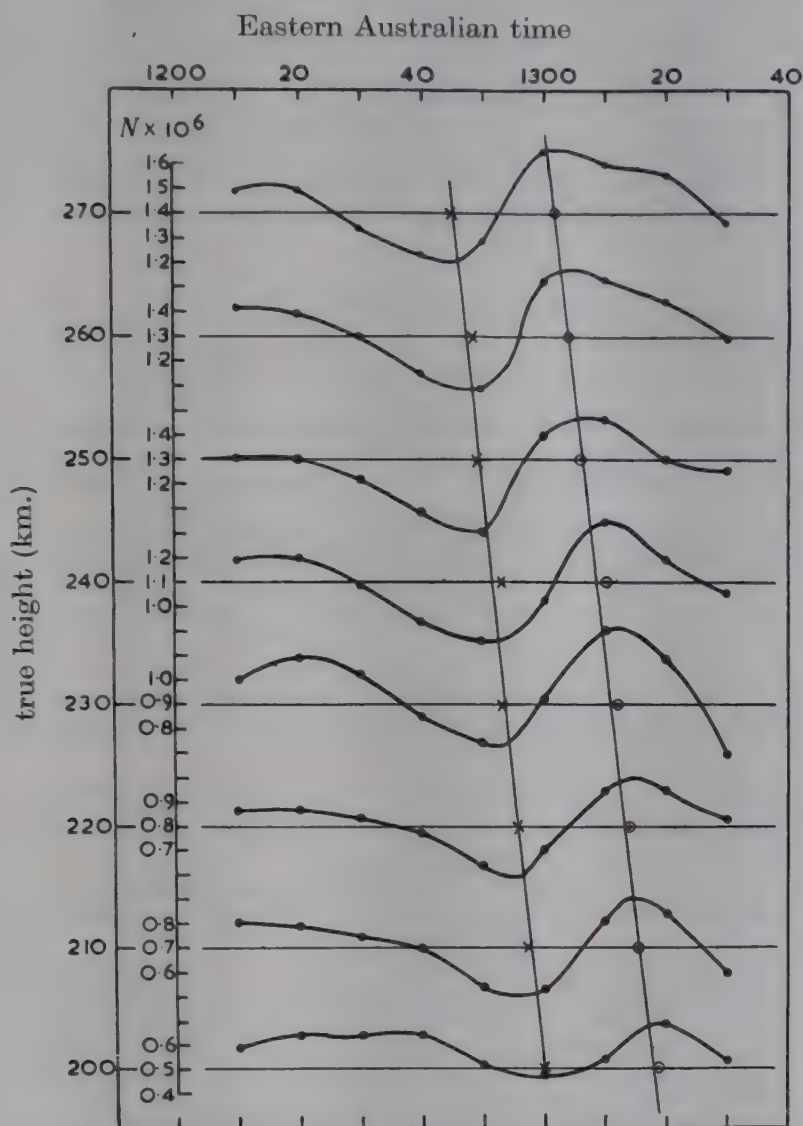


FIGURE 4. Ionization changes at various heights. Brisbane, 5 July 1948.

These curves have then been used to find the change in ionization with time at a series of heights. The results for each 10 km. of height from 200 to 270 km. are shown in figure 4. The ionization scale is adjusted just to cover the range of variation at each height. The diagram shows very clearly that the ionization undergoes an oscillation and that the oscillation shows a phase change with height, occurring later at lower heights. If the times of minimum and maximum ionization are indicated for each height, it will be seen that the downward progression occurs at a sensibly uniform rate. Examination of the Canberra $h'f$ records in the same way gave results very similar to those shown in figure 4, thus implying that the apparent downward velocity is the same at Canberra and Brisbane, and that the apparent horizontal velocity is sensibly constant over the range of heights shown in figure 4. It will be obvious, also, from both figures 3 and 4 that there is no marked change in ionization throughout the region but only a temporary redistribution.

It should be noted that examination of figures 3 and 4 alone would suggest that the disturbance is travelling vertically downwards. Since, however, in this case it has travelled horizontally at least 900 km., it appears more probable that the fundamental movement is the horizontal one, and that the apparent vertical movement is due to a change of phase with height (Martyn 1950).

Wells, Watts & George (1946), using a 'panoramic' ionosphere recorder at Washington, D.C., U.S.A., have observed 'ripples' proceeding down the $h'f$ curves. They have ascribed them to streams of particles descending vertically. It will be seen, however, that in some cases at least the effects observed may be due to disturbances of the type being discussed.

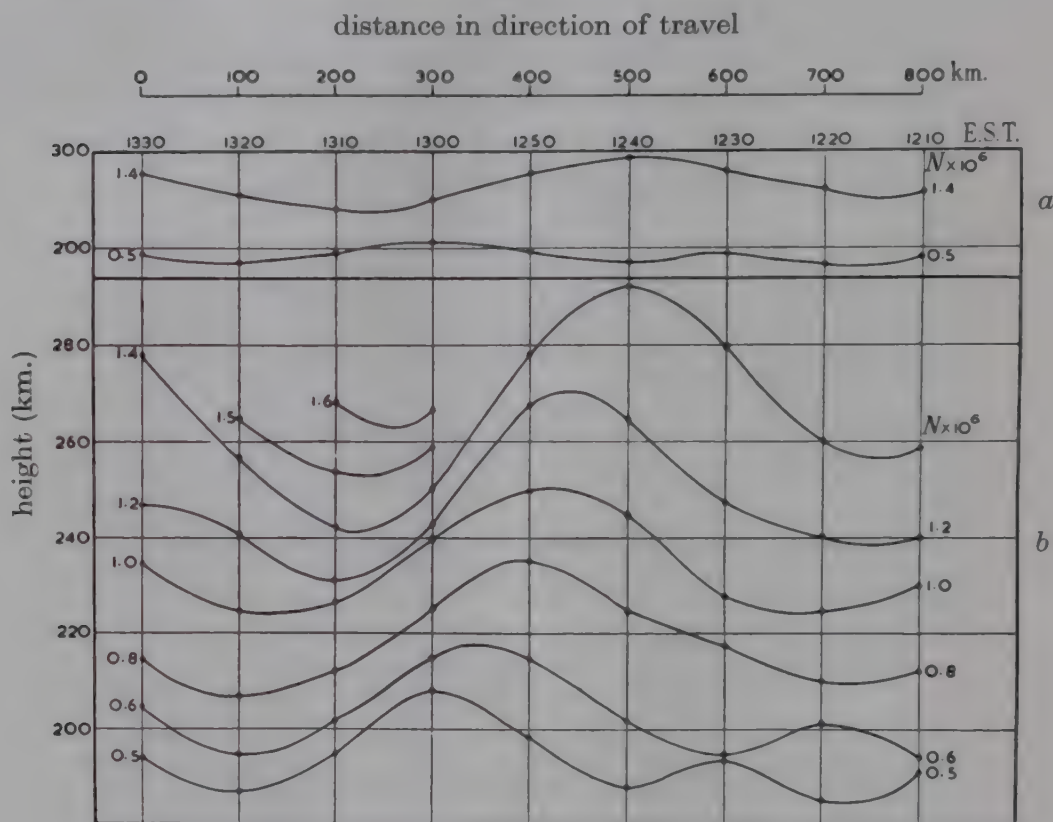


FIGURE 5. Distribution of ionization in direction of travel.

Figure 4 shows that this particular disturbance is quasi-periodic in form, consisting of one complete cycle with a period of approximately 50 min. Taking the velocity as 10 km./min., this gives a horizontal wave-length of 500 km.

It will be seen also that the apparent downward velocity is approximately 5 km./min. It is not possible to determine the full extent of the disturbance vertically owing to the absence of reflexions from beyond the height of maximum ionization and to the absence of F region echoes on lower frequencies due to E -region absorption and reflexion.

It is not possible, either, from this single case to determine the extent of the disturbance perpendicular to the direction of travel. The uniformity of the variations at the extremities of the three-point system suggests that the disturbances usually have horizontal travelling fronts at least 50 km. long, and examination of some large disturbances on the records from Brisbane, Sydney and Canberra when

the direction of movement was transverse to the lines joining these stations indicates that the fronts may extend for some hundreds of kilometres.

These characteristics suggest that the fundamental disturbance is not a change in ionization, but a travelling disturbance of a pressure type which produces a corresponding redistribution of ionization as it passes a fixed point.

The fluctuation of the ionization about the mean is, in this case, of the order of $\pm 25\%$. Figure 5 shows the distribution of the ionization in a vertical plane parallel to the direction of travel. Figure 5*a* is drawn to scale, but in 5*b* the vertical scale is exaggerated to show the variations more clearly. The changes in the concentration of ionization and consequently in the vertical ionization gradient are very obvious.

4. VARIATION IN CHARACTERISTICS OF DISTURBANCES

4.1. *General. Methods of observation and reduction of records*

The previous section has dealt with one of the larger disturbances where the variations are sufficiently great for the changes to be followed on a standard ionosphere recorder. Owing to the greater resolution and accuracy in virtual height measurement and the improved continuity in time of observation, the three-point system is capable of recording a much greater number of travelling disturbances of smaller magnitude.

Records have been taken daily at R_1 , usually over the period 0900 to 1600, and for most days also at R_3 for a shorter period. Each record is examined for observable changes in the echoes from all three transmitters. All such cases are listed, and the directions of horizontal progression and the apparent horizontal velocities are deduced by a graphical method. The virtual heights at intervals of 1 min. are also read off and plotted for examination of disturbance periods and magnitude.

Phase-path counts have also been made for each minute on a number of days where records are good. They provide evidence of even smaller variations than are observable as virtual height changes, and even these smaller disturbances have shown time differences of occurrence at the spaced points.

Where large and clearly defined disturbances appear on the Sydney records examination has been made of the $h'f$ records for Canberra and Brisbane. The virtual heights at 5.8 Mc./sec. are read off from the $h'f$ records and plotted as in figure 2. It will be noted that, although the $h'f$ records are taken only at 10 min. intervals, it is possible to make an interpolation from examination of the $h'f$ records. For example, in figure 3 it will be seen that the virtual height peak on the ordinary ray occurred at a frequency above 5.8 Mc./sec. at 1250 and below that frequency at 1300, and its virtual height increased from 330 to 370 km. between those times. The interpolation should, therefore, show a sharp peak.

In all cases where a disturbance could be identified at the more distant points, the time differences were consistent with a sensibly uniform progression having the direction and velocity indicated by the three-point system. It has, therefore, been assumed that all the disturbances detected on the three-point system are of the same general type and the large number recorded by this means has been used for the following statistical examination without discrimination in terms of magnitude.

4.2. Periodicity

The disturbance discussed above does not appear to last for more than one clear cycle, but it may be regarded as having a quasi-period of the order of 1 hr. Disturbances of lesser magnitude frequently show repetitive cycles more clearly. A period of 1 hr. is about the maximum observed and is seen relatively infrequently, generally with large magnitude. This may be partly because the variation would not be so easily detected if its magnitude is small. Shorter periods are more common and range down to an apparent minimum of the order of 10 min. The periodicity usually remains of the same order throughout the day but may differ considerably from day to day.

4.3. Magnitude

The magnitude of the disturbance appears, in general, to decrease with decrease of the period. It should be noted, however, that the change in virtual height for a given change in ionization will depend on the mean ionization gradient at the point of reflexion. The sensitivity of the virtual height observations for detecting disturbances will, therefore, depend on the mean gradient at the reflexion point. This can best be determined from $h'f$ records. On ionospherically quiet days, when no appreciable changes appear on the virtual height record, phase-path observations still show variations, generally with the shorter periods of the order of 10 to 15 min.

The magnitude usually remains of the same order throughout any one day, and there are typically quiet days and typically disturbed days.

It is not proposed to deal with these aspects in any greater detail in the present paper.

4.4. Direction and horizontal velocity

A typical record of a small disturbance is shown in figure 6a. In this case both the ordinary and extraordinary rays are present for each of the three transmitters. The features for which it is possible to observe time differences of occurrence are

- (a) virtual height peaks;
- (b) virtual height dips; and
- (c) cross-over point of 'o' and 'x' rays. (In cases (a) and (b) time differences may be observed on either 'o' or 'x' rays or both.)

(d) There is also another phenomenon which is of frequent occurrence and is capable of more precise location in time. From its appearance on the records we have called it a 'Z'. A clear example is shown in figure 6b.

A 'Z' is a manifestation of a particular phase of a disturbance of the type under discussion. It is considered to be an essentially 'optical' effect involving refraction and focusing, but the details must be left for discussion elsewhere. It will be noted that there is a marked increase of intensity during the 'Z', particularly at the start thus making it easy of observation photographically.

Figure 6b shows the pattern produced by the reflected rays from the three transmitters as observed at R_3 . Both the 'o' and the 'x' rays are visible at the left on the traces marked T_1 and T_3 but, as the transmission from T_2 is weaker, only the 'o' reflexion is recorded. It will be noticed that the 'Z's' are of slightly different

duration in this case. Under such conditions the time of occurrence of the 'Z' is taken as half-way between the beginning and the end of the discontinuous portion.

The three-point system records have now been taken on the majority of days for over a year starting from April 1948.

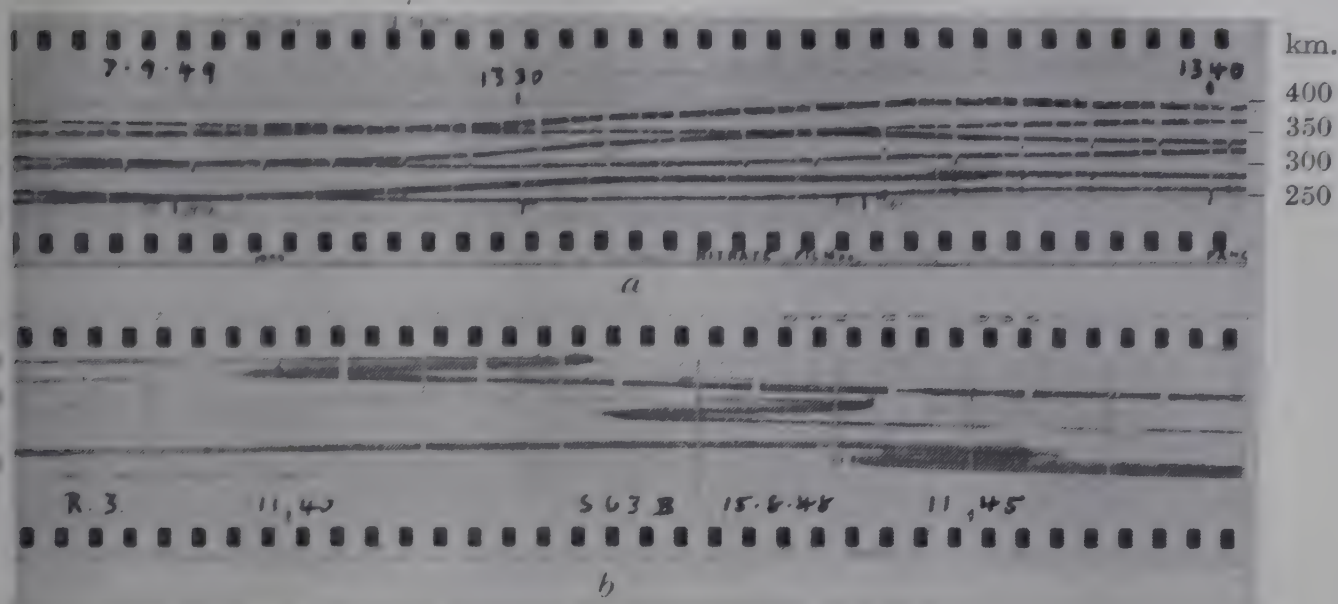


FIGURE 6. *a*, virtual height variations for 3 transmitters.
b, occurrence of 'Z's' on virtual height records.

The direction of horizontal progression and the apparent horizontal speed in that direction will be called 'direction' and 'speed' for the sake of brevity.

The accuracy of determination of the time of occurrence of an observable feature of a disturbance at all three stations is seldom greater than $\frac{1}{4}$ min. Times of occurrence were therefore recorded to the nearest $\frac{1}{4}$ min. in the majority of cases and occasionally, if the accuracy were less, to the nearest $\frac{1}{2}$ min. If this degree of accuracy were not possible, no record was made. The maximum errors due to this limitation may be as much as $\pm 10^\circ$ in bearing and ± 3 km. min. in velocity for the observations based on points 1, 2 and 3 and half this for points 1, 2 and 4.

In addition, the method of obtaining the directions and speeds tacitly assumes that the movements are uniform along a front of length greater than the extent of the observing system and that the effects of small disturbances local to one or two stations will not affect the times of occurrence of the larger disturbances. These assumptions are considered justified by the consistency of observations made at the two points R_1 and R_3 , as the discrepancies occurring are normally within the limits of accuracy mentioned above. Further confirmation was given by the consistency of the time differences when four transmitters were operated at points 1, 2, 3 and 5 for a period of 3 weeks.

For the seasonal analysis all disturbances which could be detected with adequate accuracy in time have been used without selection, as no form of discrimination appeared yet to be justified. The only eliminations which have been made were on a few days when more than ten disturbances were recorded with no significant

change in bearing or speed, or where more than one indication was recorded for a single disturbance (e.g. Z's on both 'o' and 'x' rays).

In figure 7 are shown the directions of movement for all disturbances recorded over the period of four months from August to November 1948. The directions are plotted as deduced from the time differences, except that where several observations on one day give the same direction, the points are plotted in adjacent positions. Records of several disturbances were obtained on most days over this period. A direction of 30° means that the disturbance is travelling towards 30° E of N.

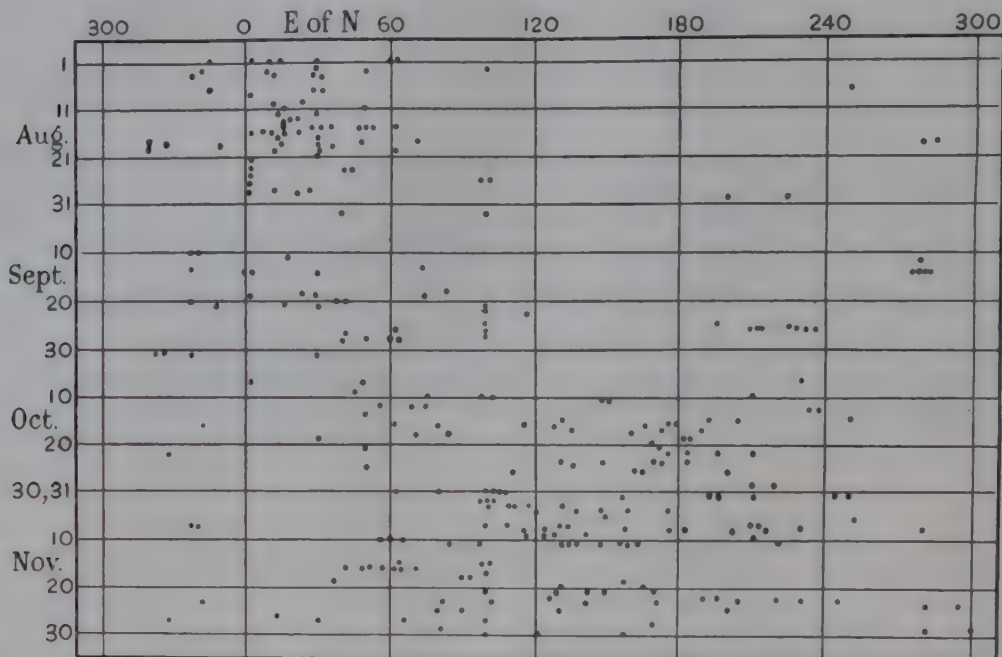


FIGURE 7. Plot of all directions of movement observed. August to November 1948.

It will be seen that in August there is a compact grouping of directions between 0 and 50° . In September the grouping is less compact, but it still has approximately the same median direction. It will be noted also that there is a complete lack of points between 120 and 180° during these months.

From 10 October, however, there is a marked change in both the median direction and the scatter. The points are now most thickly grouped between 120 and 180° and there are very few between 0 and 50° .

It is evident, therefore, that a definite seasonal change has taken place quite abruptly in the few days before 10 October. This change has been observed again in 1949 when it occurred even more abruptly between the daytime observations of 3 October and those of 4 October. There is also evidence that in 1947 it occurred between 7 October and 21 October.

Polar diagrams enable the plotting of speeds as well as directions. Figure 8 shows such a plot for the month of July 1948. The compact grouping will again be noticed. It will be observed also that the speeds in the great majority of cases lie between 5 and 10 km./min. These features are characteristic of the winter months April to September. The probable error ellipse has been calculated by the method of Chapman & Bartels (1940). A median direction obtained by taking the medians of the west-south and east-west components is also indicated by the large

cross, and it will be seen to give almost the same direction as the centre of the ellipse. Figure 9 shows the November values plotted in the same way. It will be seen that the median direction has shifted through approximately 90° , and that

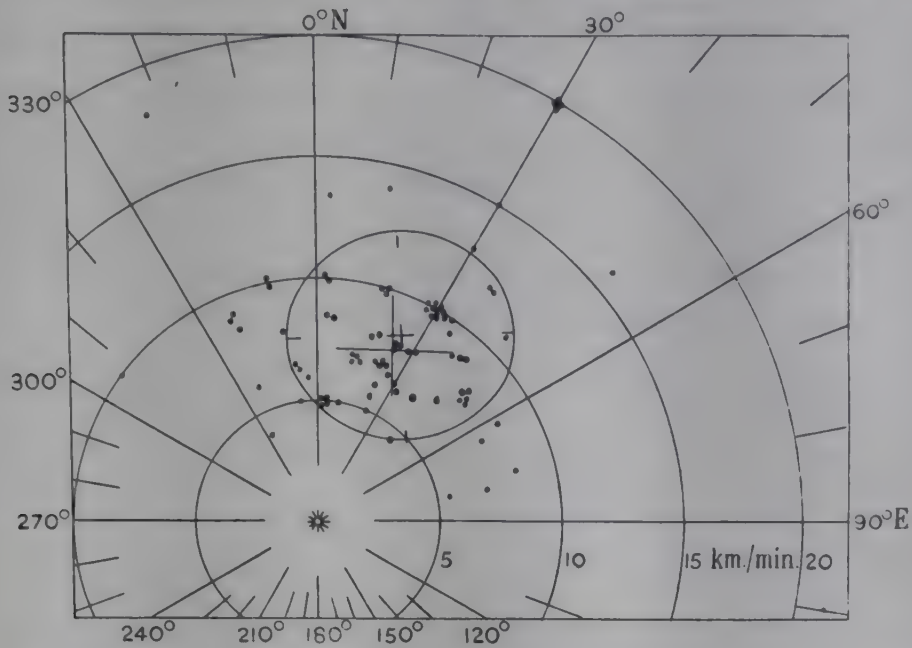


FIGURE 8. Observed velocities (directions and speeds) of horizontal movement. July 1948.

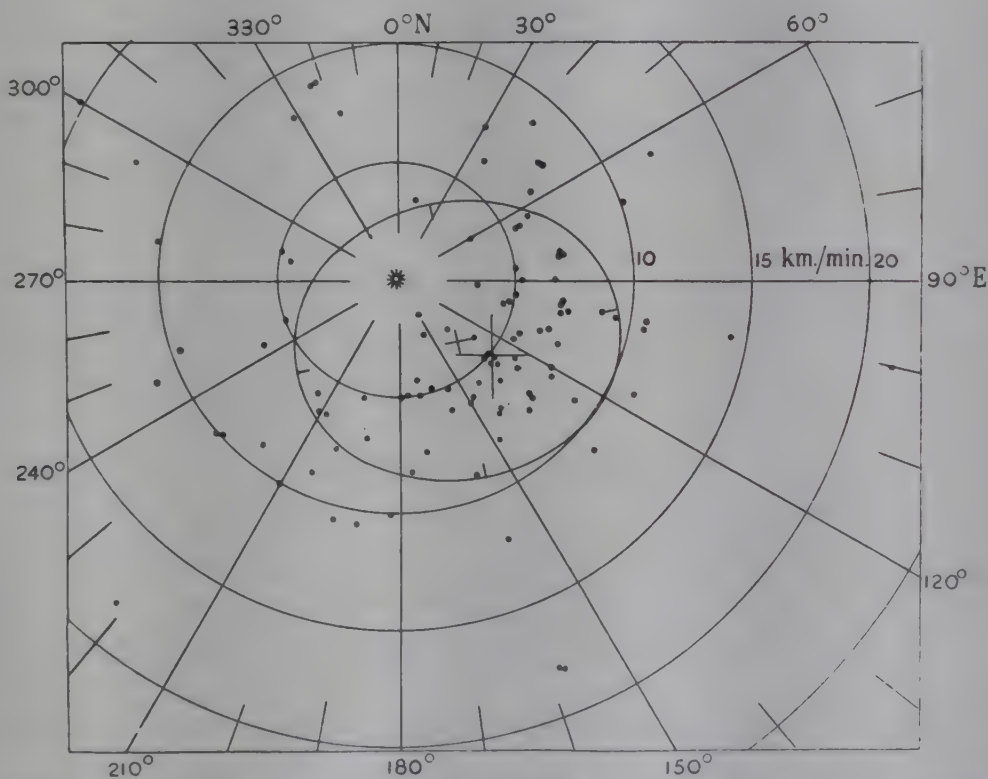


FIGURE 9. Observed velocities (directions and speeds) of horizontal movement. November 1948.

the scatter of directions is much greater. The median direction is again not greatly different from that given by the centre of the probable error ellipse. It will be noted, however, that in this case the ellipse encloses the origin. This would imply that the determination of direction and speed is not statistically reliable. From what follows, however, there seems little doubt that the median values of direction

have a physical significance. Since the speed appears sensibly independent of direction, the median of all values of velocity has been used to give a monthly median rate of movement. This rate has been used with the monthly median directions to show the seasonal variation in figure 10. It will be seen that there are two distinct groups of directions, the winter group from April to September inclusive, and the summer group from November to February. As shown in figure 7, the change from winter direction to summer direction took place suddenly just before 10 October, so that the October mean is not much affected by values prior to that date. Although the medians for March suggest a gradual return to winter conditions, the reverse change actually took place equally abruptly between the daytime observations of 16 March and of 17 March.

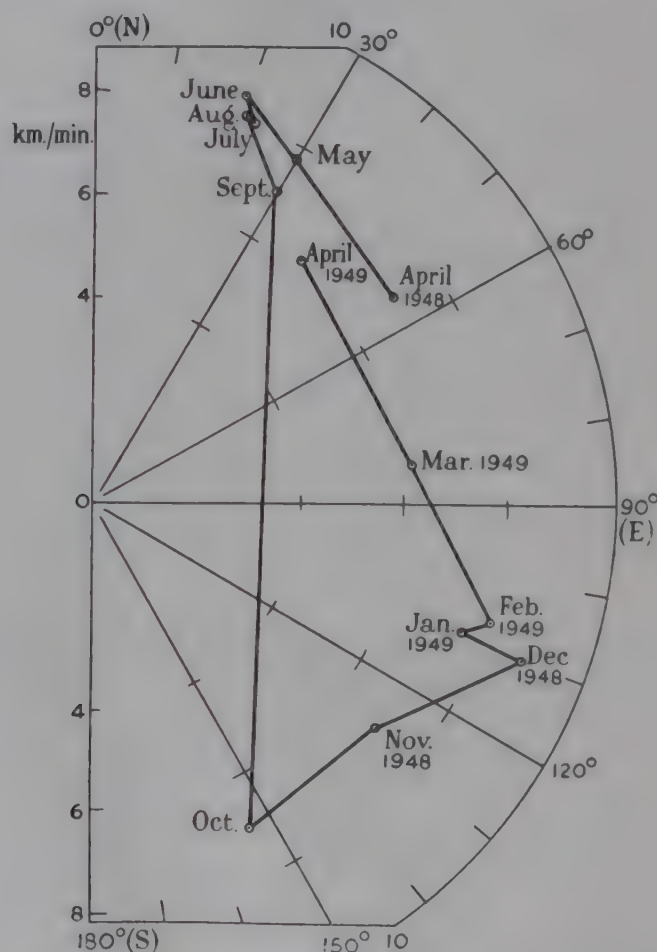


FIGURE 10. Monthly median directions and rates of movement. April 1948 to April 1949.

4.5. *Time difference in effect on the 'ordinary' and 'extraordinary' ray*

The observations of horizontal direction have, for the most part, been made on the 'ordinary' ray, as this is usually more strongly reflected. Observations on the 'extraordinary' ray do, however, give similar values for direction and speed where both rays can be observed. The delay between occurrence of changes in the ordinary and in the extraordinary ray is clearly evident in figure 3. It has been pointed out by several investigators, most recently by Millington (1949), that the two rays diverge horizontally in the plane of the earth's magnetic field as well as being reflected at different heights. The time difference in appearance of effects

on the two rays will therefore be influenced by both the horizontal and the vertical components of progression of a disturbance. In the winter months the direction of horizontal movement is approximately along the horizontal magnetic field at Sydney and the time differences for the vertical and horizontal components of progression are additive resulting in the marked 'splits' and 'cross-overs' observed. In the summer months, when the directions of movement are approximately at right angles to the field, the horizontal movement produces very little time difference and the effects on the two rays appear with a smaller time difference.

4.6. *E*-region movements

The records frequently show the occurrence of reflexions from the *E* region which persist for periods ranging from a few minutes to a few hours. It is possible to observe time difference of appearance and disappearance of the *E* reflexions and also, at times, of the disappearance and reappearance of *F* reflexions due to *E*-region 'blanketing'.

While the *E* reflexions usually have approximately similar durations at the three reflexion points, the time differences for appearance and disappearance are generally not in close agreement. From the approximate mean directions and rate of movement deduced from such observations it does, however, appear that the *E* movements are generally not in the same direction as the *F* movements occurring at the same time and the rate of horizontal movement is somewhat slower. This is in agreement with the observations of other workers, the most recent being Mitra (1949). Where *E* and *F* reflexions are observable at the same time, the *E* ionization does not appear to have any appreciable effect on the *F*-region disturbances.

5. CONCLUSIONS

The experimental evidence is considered to have established the following facts:

In south-eastern Australia quasi-periodic disturbances in the *F* region of the ionosphere with periods ranging from 10 to 60 min. are of frequent occurrence in the daytime.

The disturbances, as observed by radio methods, are evident as temporary changes in the vertical distribution of the ionization in the region.

Indication of horizontal movement is almost always present, the only definite exceptions observed being the short-period 'fade-outs' which are usually associated with solar flares. These appear simultaneously at the spaced points.

The disturbances progress horizontally as extended fronts usually at least 40 km. in length. The larger ones can travel at least 900 km. without major change in characteristics.

The horizontal directions of movement are not random but on most days are grouped about a mean direction, which is consistent from day to day over short periods. It has not yet been possible with the limited data available to detect any definite diurnal variation.

The directions do, however, show a marked seasonal variation, the major changes occurring abruptly at the equinoxes. In winter the mean direction of travel is approximately 20° E of N and in summer it is approximately 110° E of N.

The rate of horizontal travel ranges from 5 to 10 km./min. No significant diurnal, seasonal or directional change in this rate has yet been detected.

The larger disturbances, when examined, have shown also a vertical progression downwards at a rate approximately half the horizontal rate. These movements have been observed in the region between 200 and 300 km. above the earth.

It is not proposed in this paper to deal with theoretical aspects in any detail.

From consideration of the oscillatory nature of the typical disturbances, their frontal characteristics, the rate of travel and the persistence of the form of the disturbances as they travel, it seems clear that the ionization effects are only indications of a disturbance in the atmosphere itself.

The effects observed have been interpreted by Martyn (1950) as due to a cellular atmospheric disturbance in a moving air mass, the phase change with height being produced by the action of the earth's magnetic field on the ionization.

The seasonal variations in the direction of movement suggest that they may be associated with some large-scale changes in the upper atmosphere. As long ago as 1925 it was reported by Round, Eckersley, Tremellan & Lunnon (1925) that in the northern hemisphere a sudden change in radio propagation conditions took place annually in October or November and similar effects have been observed more recently by other workers. It is interesting to note also that the directions of apparent horizontal movement observed by Beynon (1948) in Europe were predominantly from east to west, which is the reverse of the directions in Australia.

It therefore seems most important that these movements and the sudden changes should be examined on a world-wide basis.

The directions of movement may be observed in several ways in addition to that described in this paper.

Comparison of $h'f$ records spaced a few hundreds of kilometres apart should show the larger disturbances and give an indication of their movement.

Beynon (1948) observed small irregularities (not specified) appearing first on oblique transmission and later on vertical transmissions. Such observations at two suitably spaced points or on oblique transmissions from two spaced points would enable deduction of a horizontal direction and velocity.

More recently Ross & Bramley (1949), using direction finders, have observed short-period variations in bearing which appear to be due to tilts in the ionosphere. These would necessarily occur under conditions such as those shown in figure 5, particularly if the disturbance be in the form of an extended front.

Observations on receivers spaced farther apart than in this case should indicate a direction of movement of the disturbance.

The sudden change in October and March may appear as any of a number of changes in radio propagation according to the radio-frequency in use and the path traversed.

This work has been carried out as part of the activities of the Radio Research Board of the Commonwealth Scientific and Industrial Research Organization of Australia, and it is published with the consent of the Executive. I wish to acknowledge with gratitude the interest of Sir John Madsen, the Chairman, and Dr D. F.

Martyn, the Chief Scientific Officer, and the assistance of Mr J. A. Harvey and the technical staff of the Sydney Section of the Board in the design and preparation of equipment, the carrying out of observations, and the analysis of the records.

Thanks are due to the University of Sydney and, in particular, to Dr D. M. Myers, Professor of Electrical Engineering, for the provision made in his Department for carrying out the work, also to the Radiophysics Laboratory of C.S.I.R.O. for the construction of some of the equipment, and to the Commonwealth Observatory and the University of Queensland for the use of records from their variable frequency recorders.

REFERENCES

- Beynon, W. J. G. 1948 *Nature*, **162**, 887.
Chapman, S. & Bartels, J. 1940 *Geomagnetism*, **2**, chap. xix. Oxford: Clarendon Press.
Manning, L. A. 1947 *Proc. Inst. Rad. Engrs, N.Y.*, **35**, 1203.
Martyn, D. F. 1950 *Proc. Roy. Soc. A*, **201**, 216.
Millington, G. 1949 *Nature*, **163**, 213.
Mitra, S. N. 1949 *Proc. Instn. Elect. Engrs*, **96**, 441.
Munro, G. H. 1948 *Nature*, **162**, 886.
Munro, G. H. 1949 *Nature*, **163**, 812.
Pierce, J. A. & Mimno, H. R. 1940 *Phys. Rev.* **57**, 57.
Ross, W. & Bramley, E. N. 1949 *Nature*, **164**, 355.
Round, H. J., Eckersley, T. L., Tremellan, K. & Lunnon, F. C. 1925 *J. Instn. Elect. Engrs*, **63**, 933.
Wells, H. W., Watts, J. M. & George, D. E. 1946 *Phys. Rev.* **69**, 540-541.

Adhesion of solids and the effect of surface films

By J. S. McFARLANE AND D. TABOR, *Research Laboratory on the Physics and Chemistry of Rubbing Solids, Department of Physical Chemistry, University of Cambridge*

(Communicated by F. P. Bowden, F.R.S.—Received 16 January 1950)

Earlier experiments have shown that the friction between metals is due to the shearing of junctions formed by adhesion or welding at the points of intimate contact. This suggests that when the load is removed the junctions should remain and an appreciable normal force should be needed to separate the surfaces. Experiments show that with clean hard surfaces in dry air the adhesion is negligibly small. In moist air appreciable adhesion may be observed, and it is shown that this is due to the surface tension of a thin film of adsorbed water. The surface-tension forces due to thin films of liquid trapped between solid surfaces may be very large. Under certain conditions the viscosity of the liquid may also be important.

The absence of adhesion between clean hard surfaces is not due to the non-formation of metallic junctions. Experiments show that it is due to the released elastic stresses which break the junctions one by one as the load is removed. With very soft metals, such as lead or indium, where the effect of released elastic stresses is very much less important, marked adhesion is observed in air, if the surfaces are freed of grosser contaminants. This adhesion provides direct evidence for the formation of metallic junctions by a process of cold welding or pressure welding at the points of contact. If the surfaces are covered with oxide films of appreciable thickness, the amount of metallic interaction is diminished with a corresponding reduction in the adhesion. Lubricant films have a similar effect, and in general those materials which are most effective in reducing the adhesion are also most effective, as boundary lubricants, in reducing the friction.

INTRODUCTION

This paper describes an investigation into the adhesion between solid surfaces in air, under various experimental conditions. Earlier experiments have shown that the friction between solids, especially metals, is due to the shearing of junctions formed by adhesion or welding at the points of intimate contact. This suggests that when the load is removed these junctions should still remain and an appreciable normal force should be needed to separate the surfaces. It is therefore of interest to examine the normal adhesion between solids and to determine the effects of interposed films on the adhesion. Earlier work by Holm (1946), Shaw & Leavey (1930) and Jacob (1912) has shown that strong adhesion may occur between surfaces which have been completely freed of adsorbed surface films, but if the surfaces are exposed to air no adhesion is observed. Similar results have been recently described by Gwathmey, Leidheiser & Smith (1948). These observations are consistent with frictional measurements by Holm (1946), Bowden & Hughes (1939) and Bowden & Young (1949), who showed that the friction between outgassed metals may be extremely high, and is considerably reduced by the presence of contaminant films.

The work of McBain (1931) and Hardy (1936) has dealt more specifically with the adhesion due to interposed films of adhesives. Here the adhesion is determined by the strength of the adhesive itself or by the attachment of the adhesive to the solid surfaces, whichever is the smaller. Similarly, Budgett (1912), Stone (1930) and other workers have shown that marked adhesion between solid surfaces can occur

in the presence of water or other liquid films. The first part of this paper shows that the adhesion in such cases is due primarily to the surface tension of the interposed liquid. Under certain conditions the viscosity of the liquid may also be important. In the second part of the paper it is shown that marked adhesion between metals may occur in air if the metals are soft and free of grosser surface contaminants. This adhesion is a direct result of a cold-welding or pressure-welding process at the points of intimate contact. The effect of surface films on the adhesion is described, and it is shown that those films which most effectively reduce the adhesion are good boundary lubricants (see also McFarlane & Tabor 1948).

1. THE ADHESION OF HARD SURFACES: GLASS, PLATINUM AND SILVER

A number of experiments were carried out to measure the normal adhesion between surfaces of glass, platinum and silver after they had been pressed together. The surfaces were cleaned by good laboratory methods. Glass was cleaned in chromic acid and then in a flame, platinum by heating in a flame and silver by abrasion and metallurgical polishing. In all cases the adhesion in room atmosphere was negligible.

For this reason a sensitive pendulum-type apparatus was used, which measured the adhesion between a spherical bead, hanging freely on the end of a fine fibre, and a flat vertical plate. The two surfaces were brought into contact and the plate was then moved back until, under the force of the gravity component, the bead fell away. The deflexion of the bead from the vertical was therefore a measure of the adhesion. With this apparatus adhesions as low as 10^{-6} g. could be readily measured.

The results again showed that in clean, dry air the adhesion between the hard surfaces was negligible. In a humid atmosphere, however, marked adhesion was observed, particularly with glass surfaces. The adhesion depended on the humidity, but at saturation the adhesion was the same as that observed if a small drop of water was placed between the surfaces. This suggests that the observed adhesion is due to the surface tension of a thin film of water adsorbed on the glass surface. A similar suggestion was made by Stone (1930), but no quantitative account was given of the forces involved.

A simple calculation shows that for a spherical bead on a flat plate (figure 1) covered with a thin film of liquid, the force of adhesion due to surface tension is given by the relation

$$Z = 4\pi R\gamma \cos \alpha, \quad (1)$$

where Z = adhesion (dynes), R = radius of bead (cm.), γ = surface tension of liquid (dynes/cm.), and α = angle of contact between the liquid and the solid. This equation applies strictly only if the thickness of the liquid film is very small, and the angle of contact α is also small.

Thus the adhesion is independent of the thickness of the liquid film and is directly proportional to R . Experiments with glass beads of various radii fully confirmed this relation. Results obtained in a saturated atmosphere are shown in figure 2. The surface tension calculated from the gradient of the straight line shown is 67.3 dynes/cm., whereas the true surface tension of water at the same temperature

is 72.7 dynes/cm. The discrepancy may be due to a defect in the apparatus, for larger scale experiments carried out to measure the adhesion between a convex lens and a wet glass plate gave a value of γ of 72.5 dynes/cm., using equation (1).

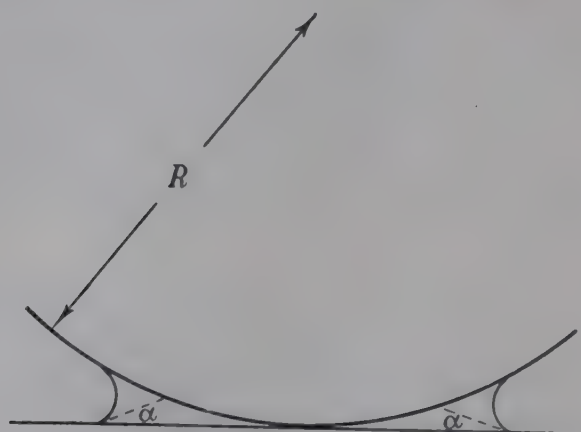


FIGURE 1

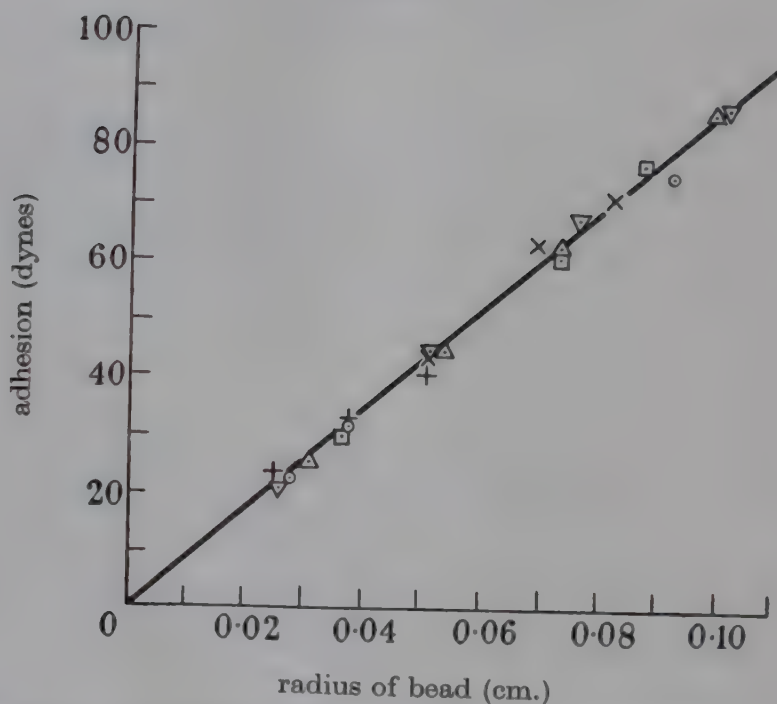


FIGURE 2. Adhesion of glass bead on flat glass surface in a saturated atmosphere.

Nevertheless, it seems clear that the adhesion measured by the pendulum apparatus is due essentially to the surface-tension forces acting in the film of water adsorbed on the surfaces. This is confirmed by experiments in which small drops of various liquids were placed between the glass surfaces. Table 1 shows the values of γ calculated from equation (1) and from the adhesions observed. Although the values are all somewhat lower than the accepted values it is clear that they are all close to and vary directly with the surface tension of the liquids.

It is of interest to note that in the above experiments the adhesion remained at its full value even when the liquid films had thinned so far by evaporation that interference colours were no longer visible. This means that liquid films well below 1 000 Å in thickness still give their 'normal' surface-tension values. We may note

incidentally that this provides a method of measuring the surface tension of minute quantities of liquid to within a few per cent. No adhesion was observed in atmospheres saturated with vapour of benzene or alcohol. This is probably because the adsorbed layers are too thin.

TABLE 1. ADHESION DUE TO THIN FILMS OF LIQUID ON GLASS SURFACES

liquid	surface tension (dynes/cm.)	
	calculated from adhesion	accepted values
water	67.3	72.7
glycerine	59	63.5
decane	22.4	25
octane	19.9	21.8

Effect of surface roughness

If the glass plate was roughened by abrading with carborundum paper, the adhesion produced in an atmosphere saturated with water vapour decreased with increasing roughness. These results are summarized in table 2. If, however, a thin film of water was applied to the roughened surfaces, high adhesions were again

TABLE 2. ADHESION OF GLASS SURFACES IN ATMOSPHERE AT 100 % HUMIDITY

surface finish	mean height of surface irregularities (Å)	adhesion as percentage of value observed with highly polished surfaces
Highly polished	c. 150	100
500 carborundum paper	1000	79
320 carborundum paper	4000	51
150 carborundum paper	100000	0

observed. Similar results were observed with platinum surfaces, but in this case, even with the smoothest polish obtainable, the adhesion was never as high as that observed between glass surfaces. This is probably because the adsorbed water film is appreciably thinner on platinum than on glass, and because the polished platinum surface is rougher than the fire-polished glass surface (see below).

Effect of humidity

Some experiments were carried out on the adhesion between smooth glass surfaces in atmospheres maintained at various degrees of humidity. The apparatus was surrounded by a lagged glass tank, and the atmosphere inside the tank was allowed to come into equilibrium with saturated solutions of various salts. The atmosphere was thoroughly mixed by a small stirrer before each adhesion measurement. With glass surfaces the adhesion reached its maximum value at about 88 % humidity; with a glass sphere on a polished platinum surface at between 93 and 100 % humidity. The adhesion results for glass are plotted in figure 3, where the adhesion is expressed as a percentage of the full adhesion observed at 100 % humidity. On

the same figure, a curve has been plotted from adsorption experiments by McHaffie & Lenher (1925) showing the thickness of the layer of water adsorbed on glass surfaces at various humidities. It is seen that the thickness of the film and the adhesion both increase rapidly at humidities greater than 90 %

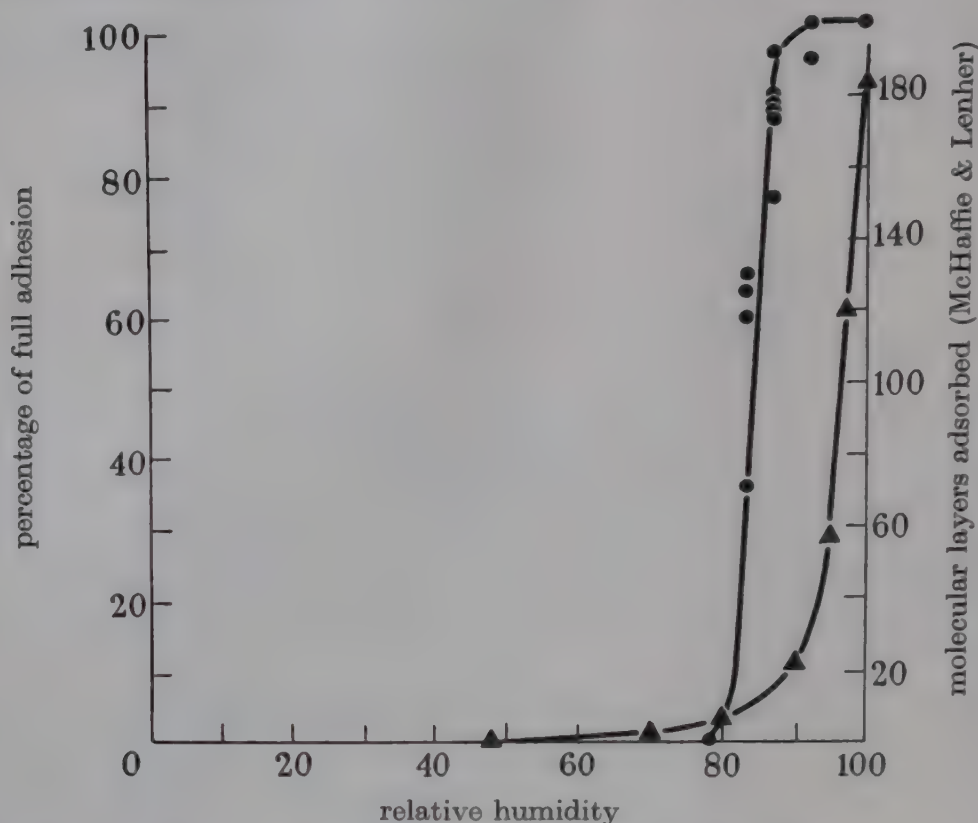


FIGURE 3. Adhesion of glass surfaces as a function of the humidity of the atmosphere. There is a close parallel between the adhesion results (curve ●) and the results of McHaffie & Lenher (1925) for the thickness of the water film adsorbed on glass surfaces (curve ▲).

The thickness of the adsorbed water film

The above experiments show that the effect of increasing the roughness of the glass plate in a saturated atmosphere is similar to the effect obtained if the surfaces are kept smooth but the humidity is reduced. This indicates that the adhesion depends on the relation between the height of the surface peaks and the thickness of the adsorbed layer of water. Simple considerations, in fact, suggest that the decrease in adhesion would occur when the height of the asperities is comparable with the thickness of the adsorbed film. Thus in a saturated atmosphere the adhesion of glass surfaces begins to fall off when the surface roughness exceeds about 1000 Å.

These results show that water films adsorbed on glass and platinum surfaces are responsible for the adhesion observed. It seems that on glass surfaces these films must be relatively thick even at humidities below saturation.

Adhesion due to surface tension and viscosity

To check the validity of equation (1) some experiments were made to measure the adhesion between a double convex glass lens (radius of curvature 5.042 cm.) resting on a flat glass plate covered with a thin film of liquid. Table 3 shows the results obtained at 20° C.

TABLE 3

liquid	angle of contact on glass	adhesion (g.)	calculated surface tension 20° C (dynes/cm.)	true* surface tension 20° C (dynes/cm.)
alcohol	0°	1.43	22.1	22.3
benzene	0°	1.88	29.1	28.9
water	0°	4.70	72.5	72.7
aniline	c. 17°	2.64	42.7	42.9
castor oil	0°	2.10	32.3	(33.0)

* From the *International Critical Tables*.

All the liquids, except castor oil, are of low viscosity, and with these the adhesion was readily determined; a difference of 0.01 g. in the load was sufficient to change the time interval between application of the load and separation of the surface from an infinite time to a fraction of a second.

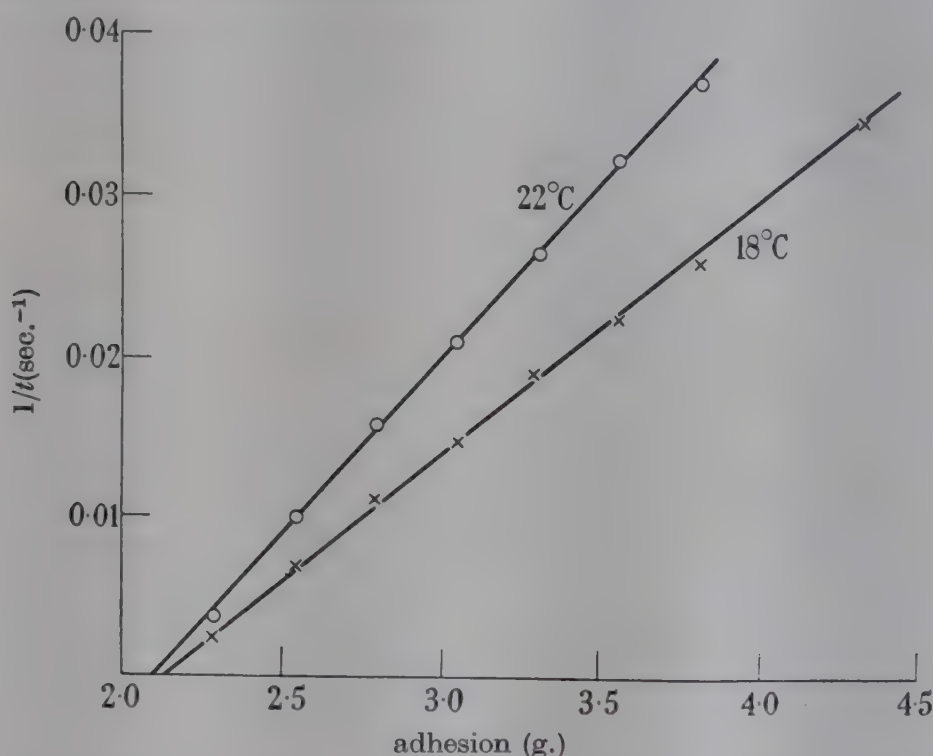


FIGURE 4. Adhesion as function of time t of separation. Glass lens resting in a pool of castor oil.

With viscous liquids such as castor oil, however, an appreciable time interval (t sec.) would elapse before the two surfaces separated, even though the load (W g.) was sufficient to overcome the surface-tension forces. If this time interval is due to the viscosity of the liquid, dimensional reasoning shows that a plot of W against $1/t$ should yield a straight line, the slope of which is inversely proportional to the viscosity. This was confirmed experimentally. Typical results for castor oil at temperatures of 18 and 22° C are shown in figure 4. The slopes of these lines are given in table 4, and it is seen that the product of the slope and the viscosity is a constant. The adhesive force at $1/t = 0$, corresponding to an infinite time of separation, gives the adhesion due to surface tension alone. It is a value obtained in this way which has been used to calculate the surface tension of castor oil in table 3.

These observations on the adhesion forces due to the surface tension and viscosity of films of liquid interposed between solid surfaces may have an important bearing

on the adhesion observed between flat surfaces such as end-gauges (see, for example, Budgett 1912; Peters & Boyd 1920; and Rolt & Barrell 1927).

TABLE 4. ADHESION DUE TO A FILM OF CASTOR OIL 0.09 CM. THICK

temp. (° C)	viscosity (centipoise)	gradient of $1/t$ against W	product of vis- cosity $\times$ gradient
18	1163	0.0160	18.6
22	834	0.0223	18.5

Suppose, for example, that a film of oil is placed between two flat circular disks, of diameter 1 in. (area 5.1 cm.^2), and suppose that the film of oil is $1000 \text{ \AA}$ thick. If the oil wets the surfaces the radius of curvature r of the meniscus at the edge of the disks is $500 \text{ \AA}$ or $5 \times 10^{-6} \text{ cm.}$ On account of the surface tension T the pressure inside the liquid film is reduced by approximately T/r . For a typical mineral oil, T is about 30 dynes/cm. , so that the reduction in pressure is $6 \times 10^6 \text{ dynes/cm.}^2$. If the film between the surfaces is continuous the adhesive force due to surface tension is $5.1 \times 6 \times 10^6 \text{ dynes} = 30 \text{ kg. (approx.)}$. The marked adhesion between flat surfaces when they are wetted by an interposed liquid was shown by Bastow & Bowden (1931). They observed that with an incomplete film of liquid the surface-tension forces are sufficient to cause a buckling of thick glass plates. These results suggest that the adhesion between Johannsen flats and similar slip-gauges may be due largely to the surface tension of an intermediate oil film. It is interesting to note that in experiments with flat surfaces in 1911, Budgett found that for highly polished steel surfaces (area 4.5 cm.^2) wetted by a film of paraffin oil, the adhesion was about 20 kg. This is of the same order as that calculated above. Budgett attributed the adhesion to the cohesive forces in the liquid film. He points out, however, that if the thickness of the film is increased the pull required to separate the surfaces diminishes rapidly. It would therefore appear that the adhesion observed is not due simply to the tensile strength of the liquid, though this will clearly set an upper limit to the adhesion (see below). On the other hand, several workers have attributed the adhesion between slip gauges to the molecular fields of the solid surfaces, acting through the liquid film (see, for example, Rolt 1929). It should, however, be noted that for dry surfaces, no adhesion is observed.

If the surfaces are entirely immersed in the liquid so that the gap between them is completely surrounded by the liquid, the surface-tension forces should be zero. This effect is, in fact, observed. If the surfaces are completely surrounded by liquid they may be separated by the smallest normal force, provided the separation is carried out very slowly. If the rate of separation is rapid the viscosity of the liquid may become a very important factor. As the surfaces are pulled apart the liquid must flow into the space between them, and if the viscosity of the liquid is appreciable the force required to separate them in a short time interval may be very large indeed. For example, a simple analysis given by Eirich & Tabor (1948) shows that if the surfaces are circular disks of radius R and the oil has a viscosity η , the time t to separate the surfaces from a separation h_1 to a separation h_2 by a force F dynes is

$$t = \frac{3\eta\pi R^4}{4F} \left(\frac{1}{h_2^2} - \frac{1}{h_1^2} \right). \quad (2)$$

Suppose the disks are 1 in. in diameter ($R = 1.27$ cm.), the oil is a typical light mineral oil of $\eta = 150$ centipoise, and the original oil film thickness is $1000 \text{ \AA}$ (i.e. $h_1 = 10^{-5}$ cm.). Then the force F required to pull the surfaces apart ($h_2 = \infty$) in a time of 10 sec. against the viscous forces (if the film is not ruptured) is 9500 kg. , i.e. nearly 10 tons. In practice, of course, the film itself will rupture at a lower load, since its tensile strength cannot reach such a high value. If dissolved air or other gases are present the oil film will rupture more readily. Nevertheless, the calculation indicates that viscous forces may produce very large apparent adhesions between parallel surfaces if the adhesive strength is measured by short-period applications of the load. The importance of viscosity in the action of adhesives has also been recently stressed by Bikerman (1947).

It is clear from these calculations and from the previous experimental results that in many cases the surface tension and viscosity of interposed liquid films may play an important part in the observed adhesion between solid surfaces. This will be particularly marked when the surfaces are very smooth and separated by a thin continuous film of liquid extending over a relatively large area.

2. THE ADHESION BETWEEN SOFT METALS

The fact that in the complete absence of liquid films normal adhesion between clean platinum, silver and glass surfaces in air is negligible may appear to argue against the formation of inter-metallic junctions. There are, however, three points to be considered. The first is that contaminant films on the surfaces may not be readily broken through when the surfaces are placed in contact, unless the deformation is very great. When sliding occurs, however, these surface films are largely broken down by the sliding process itself. The second point is the experimental difficulty in arranging for the separating force to be applied to all the junctions at once. More usually the force will tend to peel the surfaces apart, thus breaking the junctions one by one. The third, and probably the most important point, is that the solids used in the previous experiments are relatively hard. Consequently when the joining load is removed and the elastic stresses within the bulk of the bodies are released, the resultant elastic recovery may break any junctions that may have been formed (see Bowden & Leben 1940; Ernst & Merchant 1940). This view is confirmed by experiments described below using soft metals where the elastic recovery is very much less.

Effect of elastic recovery

Before describing the experimental procedure and results it is interesting to consider more quantitatively the effect of the recovered elastic stresses on the adhesion of metals. We consider the case of a hard spherical surface that is pressed into the flat surface of a softer metal. It is well known that when the load is removed the radius of curvature of the recovered indentation is greater than the radius of the sphere which produced it. This effect, which is known in Brinell hardness measurements as 'shallowing', is illustrated in figure 5. This change of shape is due to the release of elastic stresses (Tabor 1948) and may be described in

terms of Hertz's equation for the elastic deformation of spherical surfaces (Hertz 1896), viz.:

$$d = 2.22 \left[\frac{F}{2} \frac{r_1 r_2}{r_2 - r_1} \left(\frac{1}{E_1} + \frac{1}{E_2} \right) \right]^{\frac{1}{3}}, \quad (3)$$

where d = chordal diameter of indentation (cm.), F = normal force applied (dynes), r_1 = radius of ball (cm.), r_2 = radius of curvature of recovered indentation (cm.), E_1 = Young's modulus of ball (dynes/cm.²), E_2 = Young's modulus of metal (dynes/cm.²), and Poisson's ratio is assumed to be equal to 0.3. It is also assumed that the diameter of the impression, d , changes little when the load is removed.

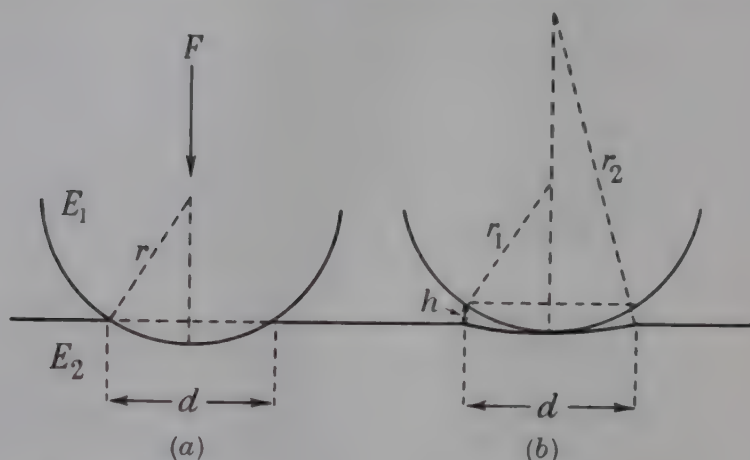


FIGURE 5. *a*, under load F dynes; *b*, load removed.

This change in shape will tend to tear apart any junctions which may have been formed when the load is initially applied. We may take the distance h (figure 5) as a measure of this tendency, and to a first approximation $h = \left(\frac{1}{r_1} - \frac{1}{r_2} \right) \frac{d^2}{8} = \left(\frac{r_2 - r_1}{r_1 r_2} \right) \frac{d^2}{8}$; substituting in equation (3) we obtain

$$h = k H d \left(\frac{1}{E_1} + \frac{1}{E_2} \right), \quad (4)$$

where h is in cm., $k = 0.53 \times 10^8$ and H = Meyer's hardness for the metal in kg. mm.². This is the ratio of the load to the projected area of the indentation and is not very different from the Brinell hardness which is the ratio of the load to the curved area of the indentation. The calculated values of h for several metals, assuming a standard value of $d = 0.2$ cm., are shown in table 5.

TABLE 5

ball	E_1 (dynes/cm. ²)	metal	E_2 (dynes/cm. ²)	H (kg/mm. ²)	h (10 ⁻⁴ cm.)
hard steel	20.9×10^{11}	mild steel	20.9×10^{11}	150	15
		copper	12.3×10^{11}	40	5.5
		cadmium	5.0×10^{11}	22	9.3
		aluminium	7.2×10^{11}	17	3.3
		tin	5.4×10^{11}	5	1.2
		lead	1.6×10^{11}	4	2.7
		indium	1×10^{11}	1	1.0

It is seen that the distance h , which would be the vertical separation between the ball and the metal at the rim of the impression in the absence of adhesion, is fifteen times greater for steel than for indium. Consequently, with steel surfaces, metallic junctions may well be broken by the elastic recovery when the load is removed, and the residual adhesion may be very small. With a softer metal such as indium, however, the amount of movement at the rim of the impression is very much smaller. In addition, the softer junctions are far more ductile and are more easily deformed. As a result the junctions may not be ruptured when the load is removed and appreciable adhesion may be observed.

Though the effect of released elastic stresses has only been discussed in detail for a spherical indenter, it is evident that similar considerations will apply in all cases where the distribution of elastic stresses under load is uneven.

Apparatus

For these experiments a very simple apparatus was used. It consists of a rigid beam, pivoted on bearings at its centre. At one end of the beam the two metal surfaces are mounted. The upper surface is usually a hard steel ball, though other surfaces were sometimes used, and the lower surface is a flat block of the softer metal. The ball is pressed on to the lower surface with a known load for a specified time. The load is then removed, and lead shot is run into the pan at the opposite end of the beam until the surfaces are pulled apart.

The surface of the steel ball was prepared by metallurgical polishing, washing in distilled water, drying in a blast of hot air, and rubbing with a pad of degreased cloth to remove any particles of polishing agent. The soft metal was cast into a block and a fresh surface prepared immediately before the adhesion experiment by cutting a thin layer off the surface with a clean steel planing tool. The surfaces so prepared are free from greasy matter, but they will be covered with thin films of oxide or adsorbed gas from the atmosphere.

Adhesion of indium

In these experiments values of the adhesion between a steel ball and a freshly cut indium surface were measured for various joining loads up to 5 kg. Two standard times of loading were used, 10 and 1000 sec., and three sizes of steel ball, $\frac{1}{8}$, $\frac{1}{4}$ and $\frac{3}{4}$ in. diameter. Figure 6 shows the results. The relation between load L and adhesion Z is seen to be linear in every case:

$$Z = \sigma L,$$

where σ may be called the coefficient of adhesion, by analogy to μ , the coefficient of friction. From figure 6 it is clear that σ is bigger for long times of loading than for short. It is also somewhat larger for small balls than for big balls, but the difference is not great. In all cases the coefficient of adhesion is about unity.

The effect of time of loading on adhesion was more thoroughly investigated by adopting a standard size of steel ball ($\frac{1}{8}$ in. diameter) and making measurements of adhesion against time of loading for each of the four loads $\frac{1}{2}$, 1, 2 and 4 kg. Figure 7 shows the curves obtained. If σ , instead of the adhesion itself, is plotted as ordinate all the points fall on a single curve (figure 8).

The effect of time of loading on σ is due simply to the creep of the indium under load. If the diameter of the impressions is measured, the adhesion is roughly proportional to the area of the impression (see figure 9), so that the adhesion/mm.² is very nearly independent of the size of the impression. It is, of course, true that the creep effect operates in the opposite direction too. That is to say, the force required to separate the surfaces is smaller if the force is applied for a long time. This effect, however, was not important under the standard experimental conditions used.

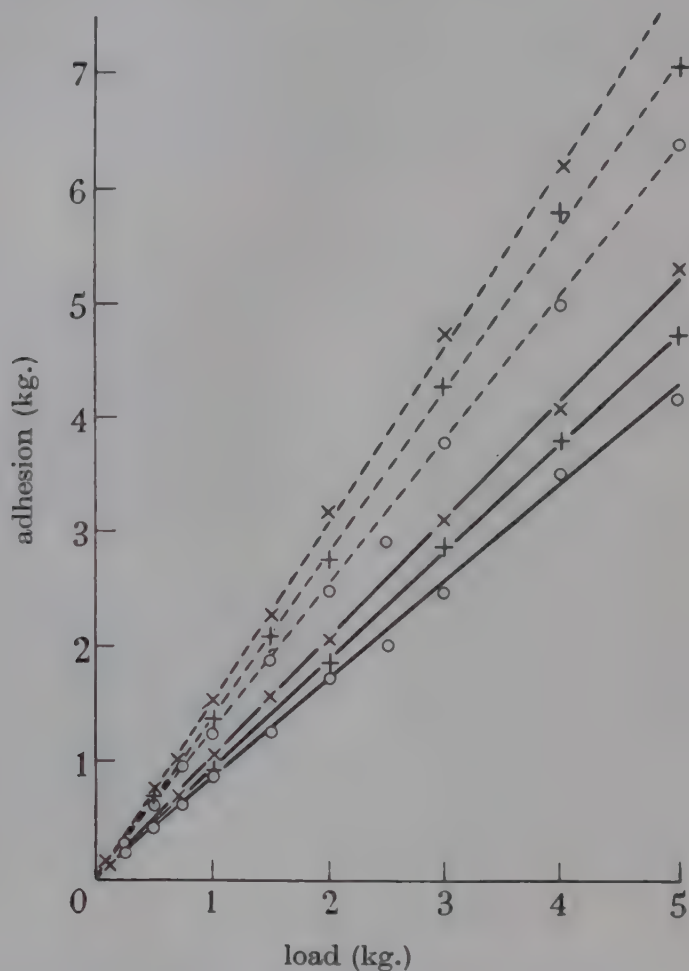


FIGURE 6. Adhesion of clean steel spheres on indium.

Diameters of spheres: $\times$, $\frac{1}{8}$ in.; $+$, $\frac{1}{4}$ in.; O , $\frac{3}{4}$ in.

Loading times: — 10 sec.; - - - 1000 sec.

When the surfaces were separated a thin film of indium was found adhering to the ball. The surface of the indentation was always very rough, and considerable 'plucking-up' of the indium had occurred. This shows that the junctions formed between the two metal surfaces are at least as strong as the indium itself. This is all the more remarkable when it is considered that the steel ball must carry an appreciable film of oxide, and that the indium must also be covered with a film of oxide or adsorbed atmospheric gases.

Measurements were also made in which the shape and relative positions of the two surfaces were interchanged. A steel flat was used as the lower surface, and a cone of indium as the upper. The adhesion between an indium cone and an indium

flat was also measured. Under all these conditions the adhesion for a given load and time of loading was substantially the same. The adhesive force in all these experiments was approximately equal to the original joining force, so that for indium on steel and indium on indium $\sigma \approx 1$.

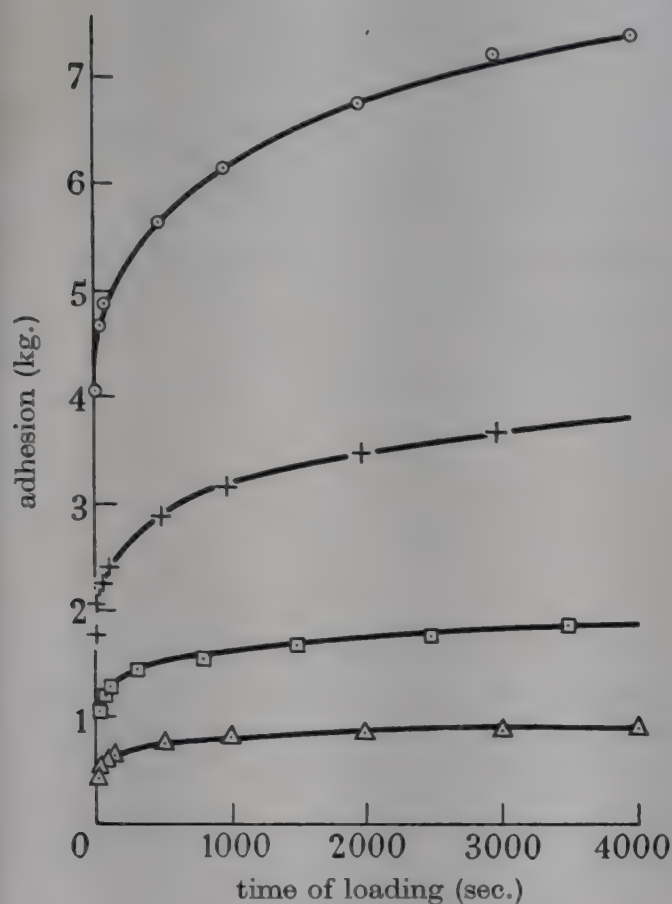


FIGURE 7. Adhesion of clean steel sphere on indium at various loads as a function of time of loading.

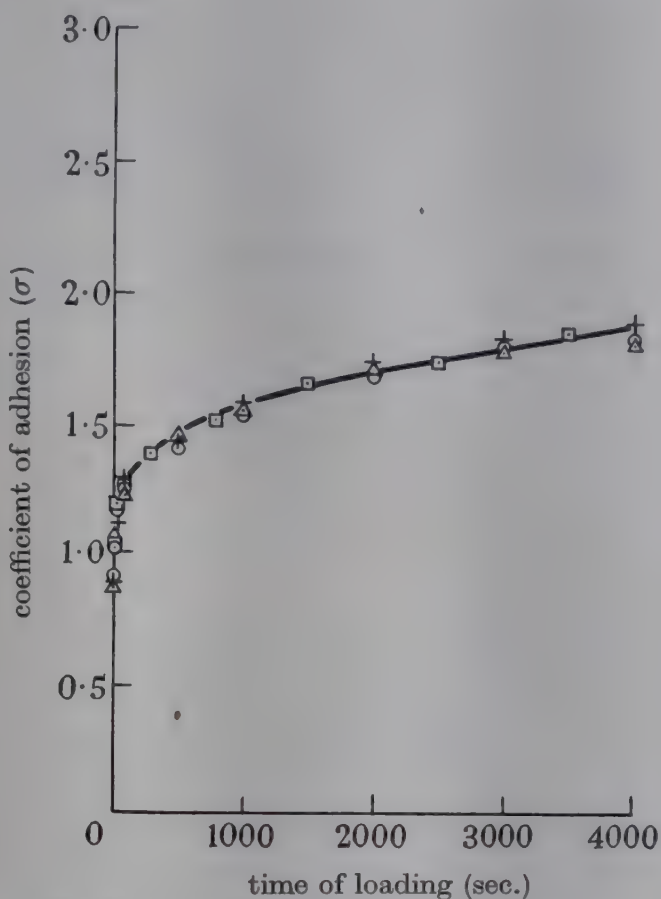


FIGURE 8. Coefficient of adhesion of clean steel sphere on indium as a function of time of loading.

Δ, 0.5 kg.; □, 1 kg.; +, 2 kg.; ○, 4 kg.

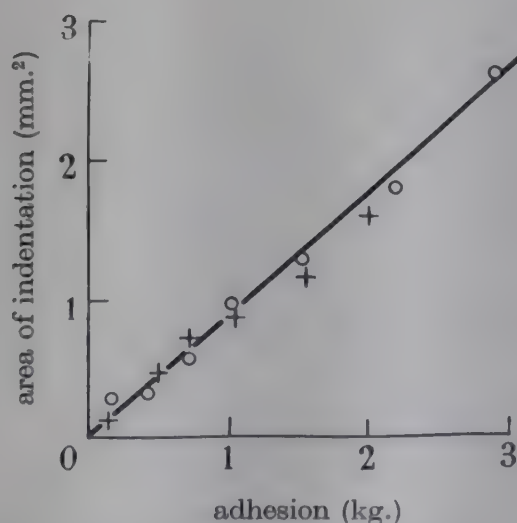


FIGURE 9. Adhesion of clean steel sphere of fixed diameter on indium. The adhesion is directly proportional to the area of the indentation for loading times of +, 10 sec. and ○, 1000 sec.

The adhesion of lead, tin, cadmium and aluminium

The adhesion between steel and lead, and steel and tin was large but not as marked as with indium. The relation between load and adhesion was found to be linear, so that the coefficient of adhesion was fairly reproducible. Cadmium gave adhesions which were large but very variable, so that the adhesion coefficient was not very consistent. With aluminium surfaces no adhesion was observed under the conditions of these experiments. The main results are tabulated in table 6, and it is seen that the coefficient of adhesion is progressively smaller for harder and harder metals. This is to be expected if the effect of elastic recovery is important.

TABLE 6. ADHESION OF STEEL ON INDIUM, LEAD, TIN, CADMIUM AND ALUMINIUM CLEAN SURFACES. LOAD 2 KG. TIME OF LOADING 1000 SEC. $\frac{1}{8}$ IN. BALL

metal surface	brinell hardness (kg./mm. ²)	coefficient of adhesion
indium	1	1.5
lead	4	0.7
tin	5	0.4
cadmium	22	0.1
aluminium	17	0.0

The effect of time of loading was investigated for lead and tin only. With these metals, as with indium, considerable creep occurred under prolonged loading. Here again the increase of adhesion for a given load produced by increasing the loading time from 100 to 10,000 sec. could be accounted for by the increase in the size of the impression alone. This was shown by the fact that the adhesion/mm.² for all these times of loading was approximately the same. The adhesion for loading times less than 100 sec., however, was anomalously small and unreproducible. The reason for this is not clear, but may be connected with the slower self-annealing of these metals compared with indium.

The specific nature of metallic adhesion

Experiments show that indium will adhere strongly to other metals as well as to steel. Some experiments were carried out, for example, with pure aluminium, which is known to carry a thin coherent film of oxide on its surface. For a cone of indium pressing on aluminium or for a sphere of aluminium penetrating indium the adhesion was very nearly the same as that observed with steel surfaces. Similar results were found for other metals such as copper and silver, and also for glass.

Attempts were also made to measure the adhesion between indium and plastics such as Perspex, polythene and polytetrafluoroethylene (teflon). Measurements were made between a hemisphere of the plastic and an indium flat, and between an indium cone and a plastic flat. With Perspex and polythene the adhesion was slight; with 'teflon' no adhesion was observed. It is evident that the adhesion of indium to other materials is specific and not indiscriminate.

Effect of oxide films on the soft metal surfaces

Adhesion experiments were carried out on surfaces of indium, lead and tin which had been exposed to air in a desiccator for periods up to 30 days. With indium there

was a steady decrease in adhesion with time of exposure. The results were as shown in figure 10. No visible change occurred in the appearance of the indium surface even after 30 days, indicating that the oxide film must be relatively thin. There was, however, a marked change in the appearance of the indentations formed in the indium surface. For surfaces which had been exposed for a prolonged period and which gave reduced adhesion, the indentations were smooth and bright, and quite unlike the distorted, plucked indentations produced in freshly cut surfaces. In addition there was much less pick-up of indium on the steel ball.

With lead and tin where the oxide is formed more rapidly and is considerably thicker the adhesion falls off more quickly. The results for lead are plotted in figure 10, and it is seen that the adhesion after 8 hr. has fallen to zero.

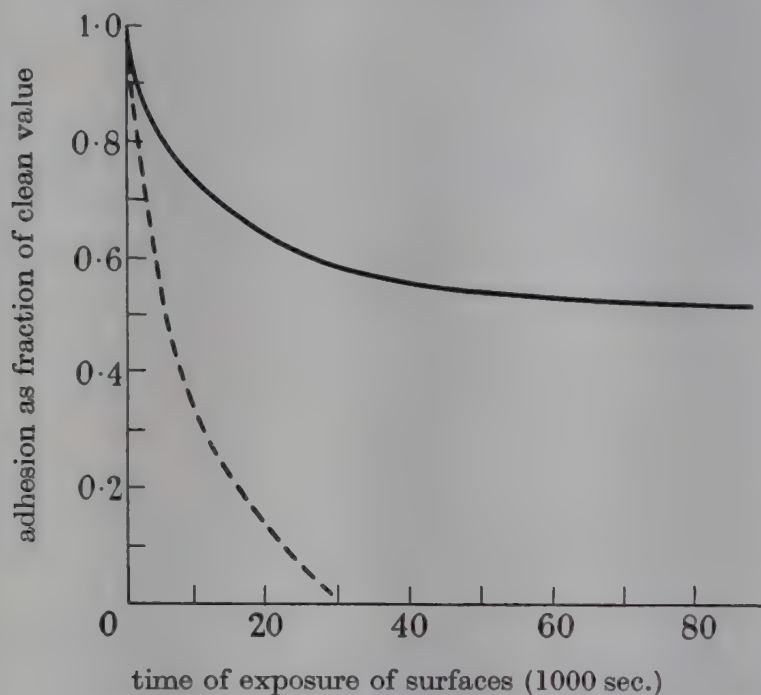


FIGURE 10. Adhesion of steel sphere on indium (—) and on lead (---) as surfaces oxidize in the atmosphere.

It is apparent that oxide films formed on the soft metal may produce a marked reduction in adhesion. This is presumably because they reduce the amount of intimate metallic contact between the surfaces. In those cases where the oxide films themselves can adhere the thick films appear to be weaker than the thin films. The effect of the oxide film in reducing the adhesion is in broad agreement with frictional observations where it is found that in general the oxide film produces a marked decrease in friction. It should, however, be borne in mind that in the frictional experiments the sliding process itself may often facilitate the breakdown of the surface film.

The effect of lubricant films

The effect of lubricant films on the adhesion was investigated. The lubricants used were a 1 % solution of lauric acid ($C_{11}H_{23}COOH$) in cetane, and medicinal paraffin oil.

(a) *Indium surfaces.* When the lauric acid solution was applied to the indium surface a monolayer of the acid was adsorbed on the metal surface and the remainder

of the solution collected as small droplets which would not wet the surface. These were carefully removed. The surface on which the adhesion experiments were carried out therefore consisted of an indium surface covered by an oriented layer of fatty acid. It was found that for small deformations the adhesion between such a surface and a clean steel ball was negligibly small. If, however, the indentation formed in the indium was sufficiently deep, marked adhesion was observed, particularly with balls of small diameter. The results suggest that the adhesion is due to a stretching and rupturing of the fatty acid monolayer. This is shown in figure 11,

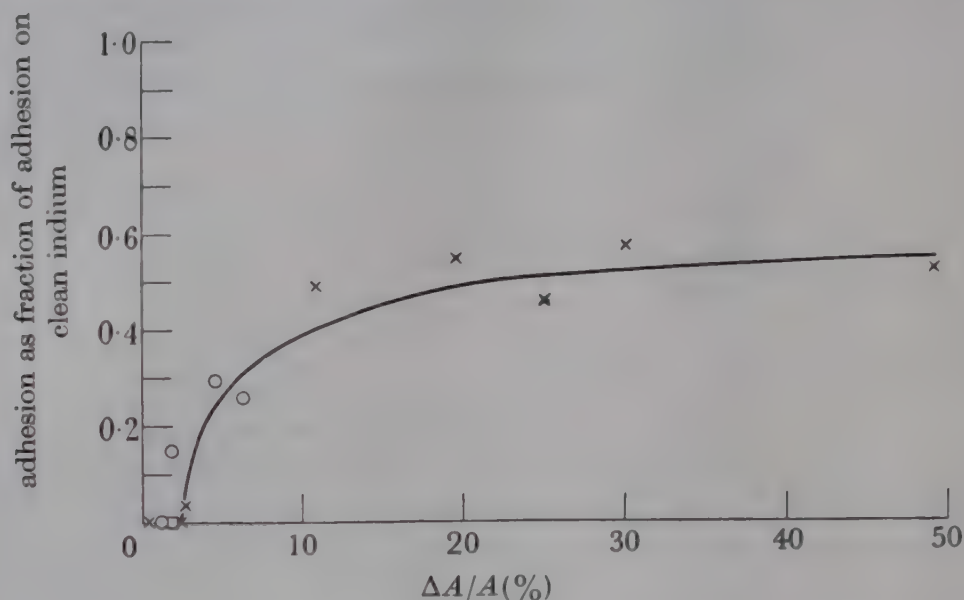


FIGURE 11. Adhesion of clean steel sphere on indium surface with a monolayer of lauric acid. Diameters of spheres: $\times$, $\frac{1}{8}$ in.; $\circ$, $\frac{1}{4}$ in.; $\square$, $\frac{3}{4}$ in.

where the adhesion is plotted against the amount $(\Delta A/A)\%$ by which the curved area of the indentation exceeds the original plane surface area. A standard time of loading of 1000 sec. was used. It is seen that for balls varying in diameter from $\frac{1}{8}$ to $\frac{3}{4}$ in. the adhesion is zero when $(\Delta A/A)$ is less than 2% but increases rapidly when $(\Delta A/A)$ exceeds this value. It would appear that if the monolayer is stretched by less than about 2% in area it is capable of preventing appreciable metallic interactions and the adhesion is negligible. If it is stretched by more than 2% the adhesion becomes large and is, indeed, much greater than would be expected if metal-to-metal contact took place only over the area ΔA . This suggests that once the monolayer is stretched by more than about 2% it undergoes very marked breakdown so that metal-to-metal contact can occur over as much as 50% of the indentation.

The view that the adhesion is due to the rupture of the monolayer is supported by two other observations. First, if a pool of the fatty acid solution is placed between the indium and the steel, no adhesion is observed even for large indentations. Secondly, if the steel ball, instead of the indium, is covered with a monolayer of the fatty acid and the ball is then pressed on to a clean indium surface no adhesion is observed, however large the indentation. Here the monolayer, which is supported by the steel surface, undergoes little deformation during the indentation process. Repeated adhesion experiments with the same ball gradually resulted in a wearing

away of the monolayer, and the adhesion gradually increased. Similar results were obtained with a 1% solution of octacosanoic acid ($C_{27}H_{55}COOH$) in cetane. These indicated that under the same experimental conditions the longer-chain acid was somewhat more effective in reducing adhesion than lauric acid.

Some interesting adhesion experiments were carried out with medicinal paraffin oil. A drop of oil was placed on a clean indium surface and the adhesion of a steel ball to this surface measured at regular intervals of time. The results are shown in figure 12, curve 1. It is seen that initially the coefficient of adhesion (σ) is 0.66 (compared with $\sigma = 1.03$ for clean indium, curve 2) and that this falls off rapidly with time, so that after an exposure of several thousand seconds it is zero. This decrease is too rapid to be accounted for in terms of the surface oxidation of the indium (compare, for example, figure 10). The results suggest that the paraffin contained small quantities of surface-active material which are slowly adsorbed on to the metal surfaces. This view is supported by some experiments with a specimen of paraffin oil that had been in contact with activated alumina for several weeks. The results are shown in figure 12, curve 3. It is seen that the initial adhesion is very much higher ($\sigma = 0.95$) and that the decrease with time is relatively slow. The very high adhesion with the purified paraffin oil indicates that the oil film is readily squeezed out or ruptured during the application of the load, enabling metal-to-metal contact to occur.

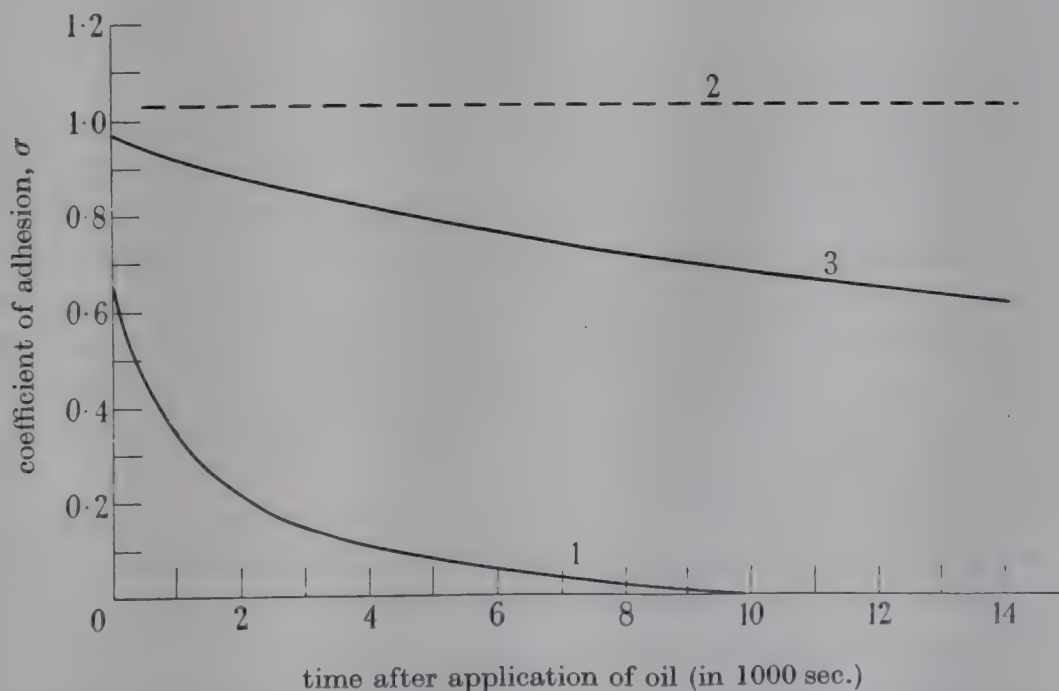


FIGURE 12. Adhesion of a clean steel sphere on an indium surface covered with a thin layer of paraffin oil. $\frac{1}{8}$ in. ball. Time of loading 10 sec.

(b) *Lead, tin and cadmium surfaces.* A few experiments were carried out with lead, tin and cadmium surfaces. In general paraffin gave a smaller adhesion than the clean surface, whilst a monolayer of lauric acid reduced the adhesion to zero for small deformations. The main results are summarized in table 7 where the adhesion is expressed as a fraction of that observed for clean surfaces.

It is seen that paraffin oil, which is a poor boundary lubricant, produces some reduction in the adhesion. With the fatty acid, which is a good boundary lubricant, even a monolayer is sufficient to reduce the adhesion to a value too small to be measured.

TABLE 7. ADHESION OF STEEL ON INDIUM, LEAD, TIN AND CADMIUM:
EFFECT OF LUBRICANT FILMS

metal	clean	adhesion		
		paraffin oil	paraffin oil purified	lauric acid monolayer*
indium	1	0.6	0.9	0
lead	1	0.15	—	0
tin	1	—	—	0
cadmium	1	—	—	0

* $\Delta A/A < 2\%$

3. DISCUSSION

The adhesion of hard metals

The experiments described in this paper show that the adhesion between solids depends on their bulk properties as well as on the nature of the surfaces themselves. With harder materials in air marked interaction may occur between the surfaces when they are pressed together, although no detectable adhesion is observed when the load is removed. This interaction is shown by frictional experiments in which the surfaces are examined after sliding has occurred. There is clear evidence that the surfaces have adhered together and that the junctions formed at these points are sheared during sliding (Moore 1948). More direct evidence has been provided by recent experiments by Mr E. Rabinowicz (unpublished) in which radioactive copper has been pressed on to clean steel. Although no normal adhesion is observed there is considerable transfer of radioactive copper to the steel surface. The absence of adhesion is largely due to the release of elastic stresses which break the junctions one by one as the load is removed.

With hard surfaces of this type strong adhesions may be observed if thin films of liquid are trapped between the surfaces. The adhesion is then due essentially to the surface tension of the liquid, and in some cases, if the rate of separation is rapid, to the viscosity of the liquid film. The adhesion will naturally be most marked if the surfaces are smooth and parallel and trap a very thin film of the liquid. It is possible that this mechanism explains the adhesion of end-gauges when they are wrung together, for it is well known that such gauges will not adhere if they are dry, nor if they are completely surrounded by liquid. Although the adhesion is due to surface tension and viscosity the upper limit to the adhesive force may be determined by the tensile strength of the trapped liquid film. If this is smaller than the forces necessary to separate the surfaces against the surface tension and viscous forces the film will rupture first. These conclusions may also have some bearing on the action of adhesives, particularly if they are capable of plastic or quasi viscous flow.

The adhesion of soft metals

With softer metals such as indium, lead and tin strong adhesions are observed in air, and there is visible evidence for the formation and rupture of metallic junctions. With indium the adhesion may be as strong as the original joining force, and this is apparently due to the absence of grosser oxide films and to the very small effects of released elastic stresses. If oxide films are allowed to grow on these surfaces or lubricant films are directly applied the adhesion is very greatly reduced.

With fatty acids a single molecular layer is sufficient to eliminate adhesion altogether. This does not necessarily mean that no adhesion occurs at all, but that the greatly reduced adhesion is now too small to withstand the released elastic stresses. Even with surfaces covered with fatty acid films large adhesions may again be observed if the surfaces are so heavily deformed that the film is broken down. This suggests that the main function of the lubricant or oxide film in reducing adhesion is to reduce the amount of intimate metallic contact. This view is supported by frictional experiments and by the great reduction observed in the amount of metallic pick-up. In some cases, however, adhesion may occur directly to the thin oxide film itself. If the adhesion with thicker oxide films is less, this would indicate that the thicker oxide layer is weaker or is less firmly attached to the underlying surface.

The effect of surface films and released elastic stresses

It is clear from these results that metallic adhesion depends primarily on two factors, the nature of the surfaces or the surface films, and the effect of released elastic stresses. Hard metals which show no appreciable adhesion in air may adhere strongly if the surface oxide film, which is normally present, is removed. This effect is enhanced if there is some sliding between the surfaces, since this tends to increase the amount of metallic contact. Thus Gwathmey, Leidheiser & Smith (1948) have shown that copper surfaces will adhere strongly if the oxide films have been removed by reduction in a stream of hydrogen at elevated temperatures. Similarly, Holm & Kirschstein (1936) found that metals outgassed in vacuum gave high frictions and often adhered together. Similar observations have been made by Bowden & Hughes (1939) and more recently by Bowden & Young (1949). For example, the latter workers found that with outgassed platinum, marked adhesion occurred after sliding even at room temperature. This result is analogous to that of Claypoole & Cook (1942), who found that clean wires could be welded by impact at room temperature.

The adhesion obtained under these conditions is due essentially to a process of cold-welding or pressure-welding at the points of intimate metallic contact. This process is used industrially in the welding of aluminium and aluminium alloys at room temperatures (see Tylecote 1948). In the experiments described in this paper it was not possible to obtain adhesion of aluminium on aluminium largely because the small degree of surface deformation was not sufficient to break up the oxide film. In the industrial process the surfaces are roughened and then very heavily deformed under pressure. This appears to rupture the oxide film and strong welds are obtained.

The importance of decreasing the elastic recovery in order to produce strong adhesion between hard metals is illustrated by many processes of pressure welding

at elevated temperatures. At these temperatures the metals are softer and the effects of the released elastic stresses when the pressure is removed will be correspondingly less. Thus Tylecote (1948) showed that the pressure welding of aluminium is greatly facilitated if the process is carried out at 200 to 660° C instead of at room temperature. Similarly Cook & Davis (see Tylecote) found that copper and copper alloys were readily welded under pressure at temperatures above 400° C, whilst Dowson (1945) has shown that silver may be easily welded by pressure alone at temperatures above 300° C. (In this case the silver-oxide film would also have decomposed.) Again the sintering of metals is greatly facilitated by being carried out at high temperatures, though here the plasticity and viscosity of the heated powders may also be important in enabling the surface tension of the metal powders to draw the granules together (Mackenzie & Shuttleworth 1949). Nevertheless, the strength of adhesion between each granule is again dependent on whether the elastic stresses can break the junctions formed.

It is sometimes suggested that high temperatures increase the rate of diffusion and the mutual solubility, and that these factors are responsible for the increased adhesion obtained in pressure welding or sintering. Although these processes play an important part it is clear that they are not the only factors. For example, very strong adhesions may occur between metals such as steel and indium at room temperature, where the mutual solubility and the rates of mutual diffusion would be negligible. In addition time has no effect on the adhesion mm.², which would not be expected if mutual solution were necessary to adhesion. Similarly, silver is a good solder for steel, and a lead-tin alloy is a good solder for zinc although the solders are not soluble in the metal concerned. Evidently strong adhesion can occur between metals which are mutually insoluble. There are however intrinsic surface properties which determine whether adhesion will, or will not occur. Thus indium adheres strongly to steel, copper or aluminium but shows no adhesion to 'teflon'. The adhesion is thus specific and not indiscriminate.

It is evident that when both the elastic recovery is marked and also the strength of the junctions is diminished by the presence of surface films the adhesion will be very small. This occurs with most common metals under normal atmospheric conditions. When either factor is largely absent some adhesion will be observed. This occurs with soft metals at room temperature, harder metals at elevated temperatures and metals from which the surface films have been removed by outgassing in a vacuum. If both factors are absent we may expect very strong adhesions. This occurs with hard metals from which the surface films have been removed and which are pressed together at elevated temperatures. It is also observed with indium at room temperature in the ordinary atmosphere.

The conditions for adhesion

It seems from these considerations that if two solid surfaces are to adhere to each other, they must be brought together over a large area of real contact without the application of great stresses since the surfaces may otherwise tend to separate elastically when the stresses are removed. The two bodies will then stick together with a force determined by the nature of the surfaces or the surface films. The

intrinsic adhesion may be great, as with indium and steel, or weak, as with indium and 'teflon'. Surface films which may modify the interaction of the metals at their interface may exert a profound influence on the strength of adhesion.

It is interesting to consider two practical examples which illustrate these principles. Adhesives are usually applied as a liquid, so that they may conform exactly to the surfaces which are to be joined without the application of stress. Later they become hard by cooling, evaporation of a solvent or by chemical change, and the strength of the joint is increased. Another example is the coherence of the grains of a polycrystalline metal, which are formed together during solidification, in close contact largely free of surface films. Here the crystals stick together so well that fracture of the metal normally occurs through the grains rather than between them.

We wish to express our thanks to Dr F. P. Bowden for valuable discussions, to the Ministry of Supply (Air) for a Research grant, to Anglo-Iranian Oil Company for special equipment and to D.S.I.R. for a maintenance grant to J. S. McF.

REFERENCES

- Bastow, S. H. & Bowden, F. P. 1931 *Proc. Roy. Soc. A*, **134**, 404.
Bikerman, J. J. 1947 *J. Coll. Sci.* **2**, 163.
Bowden, F. P. & Hughes, T. P. 1939 *Proc. Roy. Soc. A*, **172**, 263.
Bowden, F. P. & Leben, L. 1940 *Phil. Trans. A*, **239**, 1.
Bowden, F. P. & Young, J. E. 1949 *Nature* **164**, 1089.
Budgett, H. M. 1912 *Proc. Roy. Soc. A*, **86**, 25.
Claypole, W. & Cook, D. P. 1942 *J. Franklin Inst.* **233**, 453.
Dowson, A. G. 1945 *J. Inst. Met.* **71** (205).
Eirich, F. R. & Tabor, D. 1948 *Proc. Camb. Phil. Soc.* **44**, 566.
Ernst, H. & Merchant, M. E. 1940 *Friction and Surface Finish*. MIT. 76.
Gwathmey, A. T., Leidheiser, H. & Smith, G. P. 1948 *Tech. Notes Nat. Adv. Comm. Aero*, Wash., no. 1461.
Hardy, W. B. 1936 *Collected works*, Cambridge University Press.
Hertz, H. 1896 *Miscellaneous papers*, London: Macmillan & Son.
Holm, R. 1946 *Electrical contacts*, Stockholm: Hugo Gerber.
Holm, R. & Kirchstein, B. 1936 *Wiss. Veröff. Siemenskonz.* **15** (1), 122.
Jacob, C. 1912 *Ann. Phys. Lpz.*, **38**, 126.
McBain, J. W., *et al.* 1931 *Reports of the Adhesives Research Committee*, D.S.I.R.
McFarlane, J. S. & Tabor, D. 1948 *Seventh Int. Congr. Appl. Mech. London*, **4**, 31.
McHaffie, I. R. & Lenher, S. 1925 *J. Chem. Soc.* **127**, 1559.
McKenzie, J. K. & Shuttleworth, R. 1949 *Proc. Phys. Soc. B*, **62**, 833.
Moore, A. J. W. 1948 *Proc. Roy. Soc. A*, **195**, 231.
Peters, C. G. & Boyd, H. S. 1920 *Sci. Pap. U.S. Bur. Stand.* **17**, 677.
Rolt, F. H. & Barrell, H. 1927 *Proc. Roy. Soc. A*, **116**, 401.
Shaw, P. E. & Leavey, E. W. 1930 *Phil. Mag.* **10**, 809.
Stone, W. 1930 *Phil. Mag.* (7), **9**, 610.
Tabor, D. 1948 *Proc. Roy. Soc. A*, **192**, 247.
Tylecote, R. F. 1948 *Sheet and Strip Metal Users' Tech. Ass. London, Winter Conf.*

Relation between friction and adhesion

BY J. S. MCFARLANE AND D. TABOR

*Research Laboratory on the Physics and Chemistry of Rubbing Solids,
Department of Physical Chemistry, University of Cambridge*

(Communicated by F. P. Bowden, F.R.S.—Received 16 January 1950)

Simultaneous measurements have been made of the friction and adhesion of steel sliding on indium in air. The results show that both the normal and tangential stresses play a part in the deformation of the metallic junctions formed at the interface. When the surfaces are first placed in contact, a minute tangential force is required to initiate relative motion between the slider and the indium surface, since the junctions are already plastic under the applied load. As relative motion proceeds, the region of contact grows with a corresponding increase in the tangential force and the adhesive force. An upper steady state is reached where the tangential force increases more rapidly than the rate of growth of the region of contact and sliding on a macroscopic scale occurs. The detailed behaviour of the junctions during the early stages of the sliding process may be expressed quantitatively in terms of von Mises's criterion for plastic deformation under combined normal and tangential stresses, and there is good agreement between the theoretical relation and the experimental observations. The results emphasize the reality of the cold welding process which occurs at the points of intimate contact when metal surfaces are placed together. The metallic junctions so formed are responsible both for the friction and the adhesion observed. Lubricant films diminish the amount of metallic contact and so lead to a reduction in the friction and adhesion.

INTRODUCTION

Earlier work by Bowden & Leben (1939), Bowden, Moore & Tabor (1943) and by Moore (1948) has shown that there is marked interaction between metal surfaces when sliding occurs. Examination of the surface damage provides direct evidence for the formation and shearing of metallic junctions formed at the regions of intimate contact. If the surfaces are freed of oxide or other contaminant films *in vacuo* the friction reaches very high values (Holm & Kirschstein 1936; Bowden & Hughes 1939; Gwathmey, Leidheiser & Smith 1948; and Bowden & Young 1949). The surface damage is greatly increased and normal adhesion is often appreciable.

If the junctions responsible for the frictional force are also responsible for the normal adhesion it is clearly desirable to measure both quantities simultaneously and to study the correlation between them. This unfortunately cannot be done under ordinary laboratory conditions in air with harder metals, since the normal adhesion is not generally detectable. It was suggested in previous papers (McFarlane & Tabor 1948, 1950) that this is due largely to the experimental difficulty of avoiding a 'peeling' action in measuring the adhesion and to the effect of released elastic stresses. With soft metals such as lead or indium, however, large adhesions may be observed. These results suggest that with these metals friction-adhesion experiments might be carried out in air at room temperature which with other metals could be undertaken only *in vacuo* at high temperatures. This paper describes experiments on the friction and adhesion of steel sliding on indium in air at room temperature. The results show that there is a direct relation between the friction and the adhesion which may be expressed quantitatively in terms of the plastic properties of indium. The effect of lubricant films is also discussed.

APPARATUS

In these experiments an apparatus designed and used by Dr J. R. Whitehead (1950) for the measurement of friction at light loads was used. The upper surface, in the form of an $\frac{1}{8}$ in. diameter steel ball, is supported on a fine steel wire which is flexed in a vertical plane to apply the normal load. The lower surface, which is mounted on a slowly rotating turn-table, is a block of indium. When the lower surface is set in motion (at a speed of about 0.01 cm./sec.) the upper surface moves with it and flexes the suspension wire in a horizontal plane until the restoring force is equal to the friction. The deflexion of the wire, which is recorded electromagnetically, is thus a measure of the frictional force. At any stage of the sliding process, the lower surface may be stopped, the tangential force removed, and the adhesion measured. In this way the friction and the adhesion at any stage may be determined. Special precautions must, however, be taken to avoid a 'peeling' action in the measurements of the adhesive force.

FRICTION AS A FUNCTION OF LOAD

In these experiments a load of 150 g. or less was used so that the deformation of the indium was slight and the ploughing term in the frictional measurements negligibly small. For steel sliding on indium it was found that the friction rose rapidly as sliding commenced and quickly reached a high steady value. This is the maximum value attained and corresponds to that usually measured in friction experiments. The coefficient of friction at this stage depended markedly on the load. For loads above 100 g. it was fairly constant at $\mu = 2$, but at smaller loads it was much higher and for a load of 2 g. a value of $\mu = 18$ was observed. These results which are plotted in figure 1 are in marked contrast to those observed with harder metals

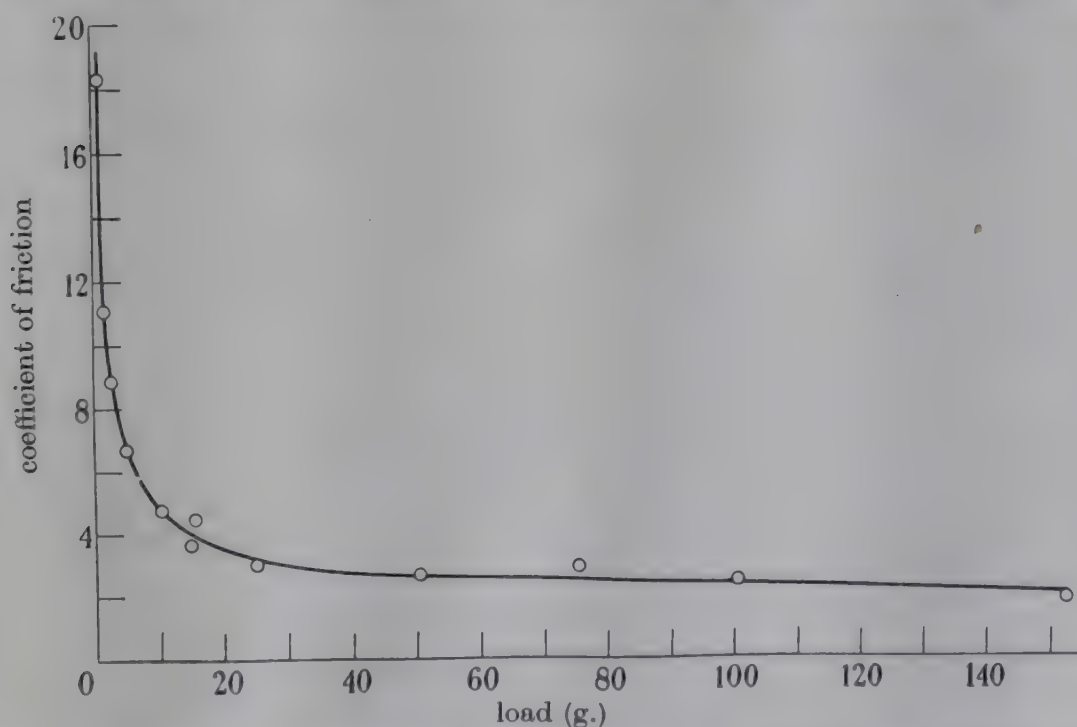


FIGURE 1. Friction of a clean steel sphere sliding on clean indium as a function of load. Above 100 g. load the coefficient of friction is approximately independent of load ($\mu \approx 2$). At small loads the coefficient rises to high values.

where the coefficient of friction is usually constant from loads of kilogrammes down to loads of milligrammes (Whitehead 1950).

FRICTION AND ADHESION

With steel on indium, if a load of, say, 15 g. was applied for a few seconds the normal adhesion was of the same order. If now the indium surface was set in motion, the frictional force on the steel ball increased rapidly to a steady upper value of about 70 g., i.e. a coefficient of friction of $\mu = 5$. If at this stage the tangential restoring force was removed, and the *normal* adhesion measured, it was found to have increased from about 15 g. to about 100 g., i.e. the coefficient of adhesion increased from about $\sigma = 1$ to about $\sigma = 7$. Thus the high frictional force is associated with a high adhesion between the surfaces.

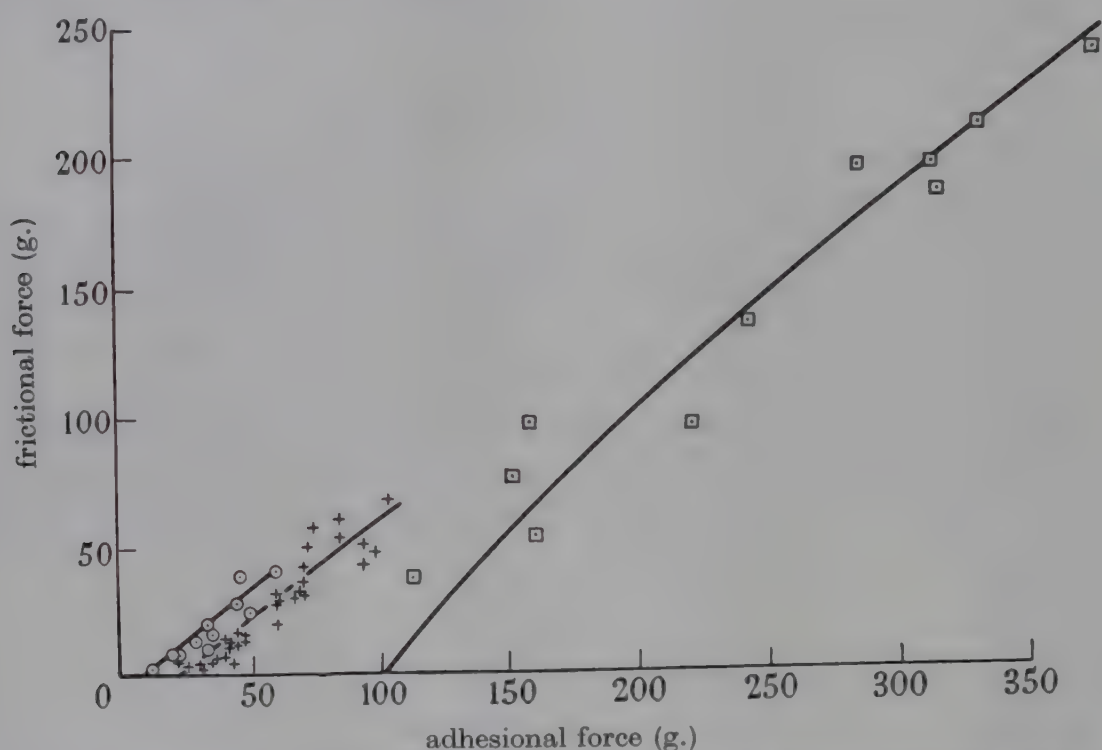


FIGURE 2. Clean steel ball on clean indium. The tangential force and adhesional force increased together as relative motion takes place between the slider and the indium. Loads: $\odot$, 5.9 g.; +, 14.5 g.; $\square$, 100 g.

It is interesting to investigate the increase in the friction and adhesion during the course of sliding. At any stage of the sliding process the lower surface could be stopped, the tangential force removed and the normal adhesion measured. The results at three different loads (5.9, 14.5 and 100 g.) are plotted in figure 2. It is seen that before sliding commences the adhesive force is approximately equal to the original load. To initiate sliding very small tangential forces are needed, but as soon as the lower surface has commenced moving there is a rapid rise both in the friction and in the adhesion. Similar results are observed when an indium sphere or cone slides on a flat steel surface. Relative motion is initiated by a very small tangential force but as the displacement continues there is a marked increase in the area of contact and in the frictional force. The upper limit is reached when sliding on a

macroscopic scale occurs, and here the frictional force may bear little relation to the original load (see also Bowden, Moore & Tabor 1943). With indium sliding on steel the large scale sliding is intermittent: with steel on indium relatively smooth. The experimental results described here are for a steel sphere sliding on indium.

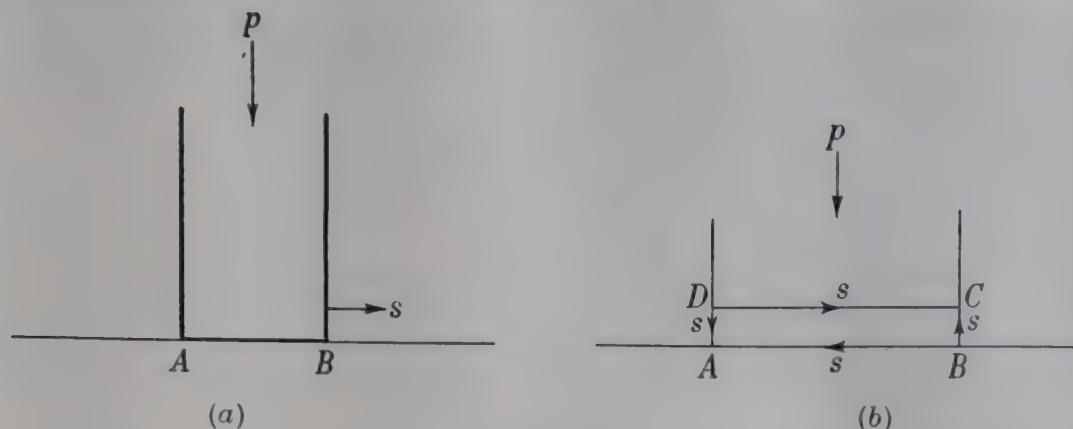


FIGURE 3

These results may be simply explained, if it is assumed that the metallic junctions responsible for the adhesion are also responsible for the friction. The frictional force is assumed to be the force required to shear the junctions. Consider for simplicity a two-dimensional model illustrated in figure 3, in which the section AB represents the region of intermetallic junctions. Suppose that intimate metallic junctions are formed as shown over the whole of the interface of area A . The normal load W produces a normal pressure $p = W/A$. When sliding occurs as the result of a tangential stress s , the metallic junctions are assumed to shear near the interface in the thin layer $ABCD$ figure 3b). In order to prevent rotation of $ABCD$ a shear stress must be assumed to act along AD and BC as shown in figure 3b. Strictly speaking there can be no shear stress in the free surfaces AD , BC , so that the shear along DC should fall to zero at D and C . The model chosen here, where the shear stress along DC is given a constant average value s , is adopted for simplicity and appears to be satisfactory. The slice $ABCD$ is thus subjected to a normal stress p and a shear stress s . Then, according to von Mises's theory of plasticity, the material in the slice $ABCD$ will flow plastically when

$$p^2 + 3s^2 = Y^2, \quad (1)$$

where Y is the elastic limit of the metal (Nadai 1931). Thus the conditions for tangential motion are determined by the combined effect of the normal and the tangential stresses. One conclusion is at once evident. If the surfaces are placed together under a normal load ($s = 0$), the indium will flow until the area of contact is sufficient to support the applied load, so that $p = Y$. Then as soon as s is given a finite value by setting the lower surface in motion, the junctions must deform, that is to say, the tangential force to initiate sliding is negligibly small. In other words, if the metal at the points of contact is plastic as a result of the normal load, only an infinitely small tangential stress is required to produce tangential flow, since the metal is already plastic. In the three-dimensional case it is not possible to solve rigorously the behaviour of the metal at the points of contact, but we may expect a general relation of the type

$$p^2 + \alpha s^2 = k^2 \quad (2)$$

at each junction. Here again, if the normal load produces plasticity at the points of contact, the tangential force to initiate sliding will be zero. This is observed. As the junctions deform, however, the metal will flow around the region of contact (figures 4 *a, b*) so that the area of contact will increase. If the normal load is kept constant, p will therefore decrease, so that if sliding is to continue and equation (2) is to be fulfilled, s will have to increase. This will lead to the steady increase in the frictional force which is observed.

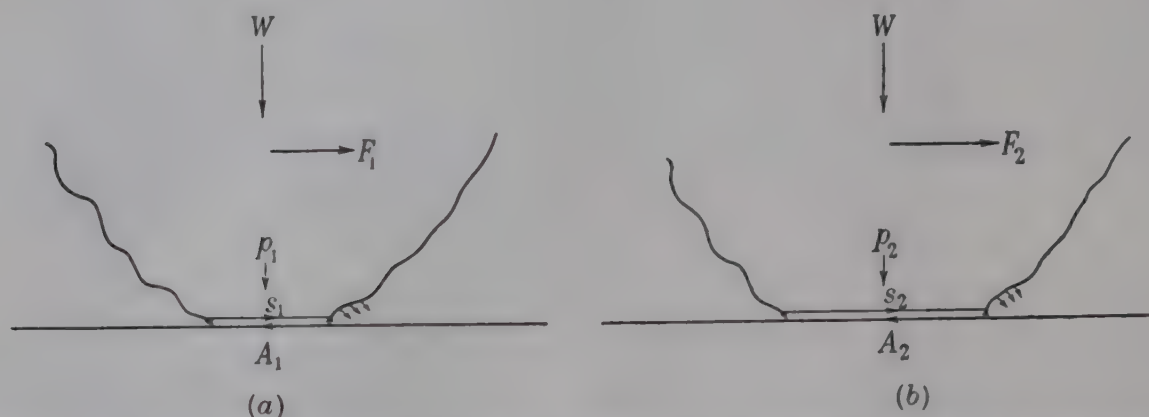


FIGURE 4. Growth of region of contact A as relative motion between the slider and the surface proceeds (*a* to *b*). Since the load W is constant, the normal pressure p decreases as the area of contact increases. Consequently the shear stress s must increase if plastic flow is to continue.

Suppose the initial load is W and the initial area over which plastic flow occurs as the result of the load is A_0 . Then $p_0 = W/A_0$. It is assumed that over the area A_0 junctions are formed which must be sheared during sliding. Before sliding occurs, $s = 0$, so that from equation (2) $p_0 = k = W/A_0$. Suppose now the normal load is removed so that $p_0 = 0$. If the junctions remain unimpaired, sliding can only be initiated by applying a finite tangential force F_0 producing a shear stress $s_0 = F_0/A_0$. The condition for flow is then given by equation (2) by putting $p_0 = 0$. Hence

$$\left. \begin{aligned} \alpha s_0^2 &= k^2, \\ \alpha &= (k/s_0)^2 = (W/F_0)^2. \end{aligned} \right\} \quad (3)$$

This ratio may be measured directly. A load of $W = 14.5$ g. was applied and then removed. The lower surface was then set in motion until slip occurred. The tangential force F_0 (as the mean of 20 readings) was 8 g. Hence

$$\alpha = (W/F_0)^2 = 3.3.$$

Equation (2) therefore becomes

$$p^2 + 3.3s^2 = (W/A_0)^2. \quad (4)$$

At any stage in the sliding process when the area of contact has spread to A , the normal pressure $p = W/A$. The tangential force is μW , where μ is the coefficient of friction at that instant, so that the shear stress $s = \mu W/A$. Substituting in equation (4) it is seen that W cancels and we are left with

$$1 + 3.3\mu^2 = (A/A_0)^2. \quad (5)$$

The ratio A/A_0 may easily be deduced from the adhesion coefficient σ . Before sliding commences it is clear that $A = A_0$. Experiments show that at this stage $\sigma_0 = 1.01$. If at any subsequent instant the adhesion coefficient has a value σ , then at this instant $A/A_0 = \sigma/\sigma_0 = \sigma/1.01$, since as has been shown in the previous paper, the adhesion is directly proportional to the area of contact. Equation (5) thus finally reduces to the relation

$$\mu^2 = 0.3\sigma^2 - 0.3. \quad (6)$$

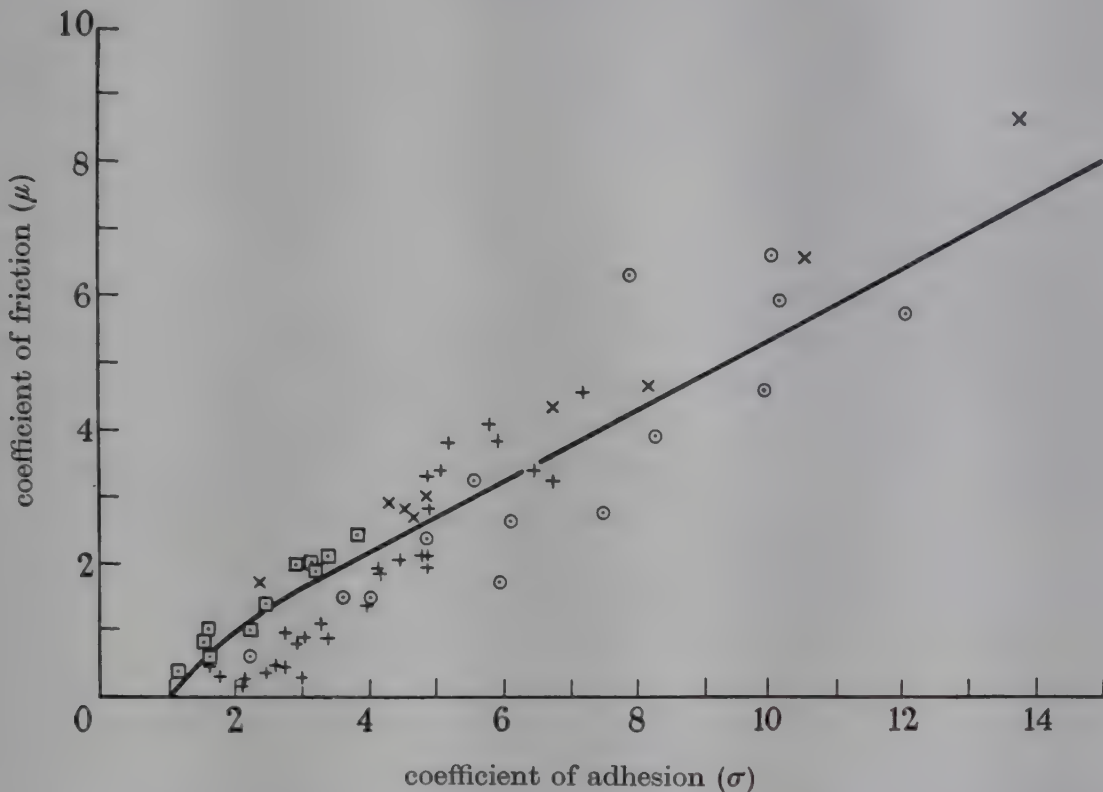


FIGURE 5. Coefficient of friction (μ) and coefficient of adhesion (σ) for steel on indium as relative motion takes place between the surfaces. Theoretical curve —. Loads: $\circ$, 5.9 g.; +, 14.5 g.; $\square$, 100 g.; $\times$, various other loads.

All the results of figure 2 have been plotted in figure 5 as μ against σ , and the thick line is the theoretical curve obtained from equation (6). In view of the experimental errors, the agreement between the experimental points and the theoretical curve is very satisfactory.* It should also be noted that this simple analysis is possible because indium does not work-harden at room temperature, so that the criteria for plastic yielding remain essentially unchanged throughout the course of the sliding process.

There are two direct observations which support the essential validity of the above treatment. The first is that when the surfaces are loaded, the smallest tangential force can initiate sliding. This is shown most sensitively by an examination of the indentation formed in the indium surface. Before sliding commences the indentation is small. If the lower surface is set in motion the friction trace suggests that the two surfaces are moving together, i.e. no relative motion is occurring between

* If the Tresca-Mohr (maximum shear-stress) criterion is adopted instead of the von Mises criterion, the same result as equation (6) is obtained.

the surfaces. This, however, is because the method is not sufficiently sensitive. If the indentation is examined microscopically it is found to have grown appreciably in size, indicating that relative motion has occurred between the surfaces before macroscopic sliding has occurred. The amount of movement involved is, of course, very small.

The second observation is the behaviour when the normal load is applied and then removed. If the lower surface is set in motion the tangential force increases. If it is kept below the critical value s_0 (see above) no tangential sliding occurs, and at any stage up to this point the *normal* adhesion and the appearance of the indentation remain unchanged. This means that there is no plastic flow of the junctions and no further increase in the area of contact until s reaches a critical value. When s reaches its critical value the junctions begin to shear, and as there is no normal load to produce a build-up of the contact area, the junctions ultimately shear completely and the friction falls to zero. In the presence of a normal load the area of contact increases very rapidly for small movements of the slider relative to the lower surface and a corresponding increase in the friction and adhesion occurs. Eventually a stage is reached where the rate of increase of the contact area is slower than the rate of increase of the tangential force. At this stage sliding on a macroscopic scale, as distinct from the minute deformation and growth of the original junctions, occurs. The friction now reaches a high value which remains relatively constant.

It is clear from these results that the flow of the junctions at the rubbing interface depends on the combined effects of the normal and tangential stresses.

OTHER METALS

Similar experiments were carried out with steel sliding on lead and tin. With these metals the friction was more or less independent of the load at a value of about $\mu = 1.2$, but at small loads there was a tendency for μ to increase. This was most marked with lead, but the effect was far less pronounced than with indium. The adhesion measurements were very irregular and irreproducible, and for this reason no further friction/adhesion experiments were carried out with these metals.

EFFECT OF LUBRICANTS

Experiments were carried out with purified paraffin oil as a lubricant between the steel and indium surfaces. For unlubricated surfaces the friction for a load of 15 g. reaches an upper value of about $\mu = 5$. In the presence of paraffin oil the motion is irregular but reaches an upper value of about $\mu = 0.9$. The adhesion is marked, of the order of 26 g., so that the coefficient of adhesion $\sigma = 1.8$. It is seen from figure 5 that this point lies on the same μ/σ curve as that obtained for clean surfaces. Thus the paraffin oil does not appear to interfere with the formation of intimate metallic junctions at the rubbing interface. It is, however, able to hinder the growth of the region over which plastic flow and adhesion occur, so that the maximum friction reached ($\mu = 0.9$) is considerably less than that observed with unlubricated surfaces

($\mu = 5$). It is probable that thin oxide layers and other contaminant films fulfil a similar function in the sliding of harder metals under ordinary laboratory conditions. Here, even in the absence of applied lubricants, the coefficient of friction rarely exceeds $\mu = 1.2$.

With a 1 % solution of lauric acid the lubrication was far more effective. The sliding was smooth and the coefficient of friction was about $\mu = 0.2$. No adhesion could be measured. Here again, this does not mean that no metallic junctions are formed, but that the amount of metallic adhesion is so small that it is unable to withstand the released elastic stresses when the load is removed. It is clear that in contrast to paraffin oil the fatty acid monolayer produces a great reduction in the amount of metallic adhesion across the interface.

DISCUSSION

The experiments on the friction and adhesion of steel sliding on indium shed further light on the mechanism of metallic friction. When the two solids are placed in contact plastic flow occurs at the points of real contact until the area is sufficient to support the applied load. At these regions metallic junctions are formed as a result of a process of cold welding. These junctions are largely responsible for the adhesion and for the friction observed. The adhesion is a measure of the tensile strength of the junctions, the friction a measure of their shear strength. When a tangential force is applied the smallest shear stress is sufficient to initiate tangential flow since the junctions are already plastic under the applied normal load. Thus at the commencement of sliding the friction is zero. The minute displacement which results leads to further flow around the regions of contact and the area of the junctions is increased. The normal load is now insufficient to produce plastic flow of the enlarged junctions and the tangential force must be increased for sliding to continue. Consequently there is a continuous increase in the area of contact with a corresponding increase in the frictional force and in the adhesion. This process continues until the rate of increase of the tangential force is greater than the increase in the size of the junctions whereupon sliding on a macroscopic scale, that is, ordinary sliding, occurs. The friction and adhesion now reach an upper value which is relatively constant. The factors determining this upper value have not been established but they are probably connected with the geometry of the surfaces and with the mechanical properties of the apparatus. In the experiments described, coefficients of friction of the order $\mu = 10$, and coefficients of adhesion of the same order have been observed at small loads.

If lubricants are applied to the surface the behaviour may be profoundly affected. With a poor boundary lubricant such as paraffin oil it would seem that the metallic interaction is not markedly affected. The friction and adhesion at the initial stages of sliding are similar to those observed with unlubricated surfaces. The main effect of the lubricant film is to hinder the growth of the metallic junctions so that the final maximum value of the friction is much lower. For example, at a load of 14.5 g. the final coefficient of friction for clean surfaces is $\mu = 5$ whilst in the presence of

paraffin oil it is of the order of $\mu = 0.9$. With fatty acids, however, a single molecular layer is sufficient to produce an enormous reduction in the amount of metal-metal interaction. The coefficient of friction is low ($\mu = 0.2$) and the normal adhesion is negligibly small. It is evident, therefore, that poor boundary lubricants limit the ultimate amount of metallic contact, whilst good boundary lubricants reduce it at all stages of the sliding process to a very low value.

With most common metals under ordinary atmospheric conditions, the coefficient of friction is generally independent of load and for unlubricated sliding has a value of the order of $\mu = 1$ (Whitehead 1950). It is interesting to consider why these metals do not show the high coefficients of friction observed with indium. There are two possible explanations. It may be suggested that the area of contact cannot be built up to the same degree because of the hardness of the metals. As the area of contact grows the normal stress decreases, so that junctions formed at an earlier stage when the normal stress was greater may be broken by the released elastic stress. Thus the final steady value of the area of contact would not show the large increase observed with softer metals. A second and more likely explanation is that it is due to the effect of surface films of oxide. These may act in a manner similar to that observed with paraffin oil on indium surfaces. They may partially reduce the initial amount of metallic contact but their main function is to inhibit the growth of the welded junctions so that the ultimate coefficient of friction is limited to a fairly constant value. This view is supported by the frictional observations of Bowden & Hughes (1939) and the more recent experiments of Bowden & Young (1949) on outgassed metals. Even with metals which are not appreciably softened by the outgassing process large coefficients of friction are obtained and marked adhesion is observed *after* a small amount of sliding has occurred. The behaviour is, in fact, similar to that observed with indium under ordinary atmospheric conditions. It should be noted, however, that with most metals at room temperature appreciable work-hardening will occur during the sliding process. This may introduce further limitations into the growth of the junctions.

These results emphasize the reality of the cold welding process when metals are placed in contact. The investigation shows that the junctions so formed are responsible both for the friction and the adhesion between metal surfaces. Surface films diminish the amount of metallic contact and the strength of the junctions with a consequent reduction in the friction and adhesion.

We wish to express our thanks to Dr F. P. Bowden for his constant interest in the work, to Dr R. Hill and Dr D. G. S. McLellan for a number of helpful discussions, to the Ministry of Supply (Air) for a Research grant, to Anglo-Iranian Oil Company for special equipment and to the D.S.I.R. for a maintenance grant to J. S. McF.

Since this paper was submitted for publication, a paper by R. C. Parker and D. Hatch (1950), *Proc. Phys. Soc.* **63**, 185, has appeared, on the friction and adhesion of indium and lead sliding on glass. Their experimental results support the general conclusions published in this paper.

REFERENCES

- Bowden, F. P. & Hughes, T. P. 1939 *Proc. Roy. Soc. A*, **172**, 263.
 Bowden, F. P. & Leben, L. 1939 *Proc. Roy. Soc. A*, **169**, 371.
 Bowden, F. P., Moore, A. J. W. & Tabor, D. 1943 *J. Appl. Phys.* **14**, 80.
 Bowden, F. P. & Young, J. E. 1949 *Nature*, **164**, 1089.
 Gwathmey, A. T., Leidheiser, H. & Smith, G. P. 1948 *Tech. Notes Nat. Adv. Comm. Aero., Wash.*, no. 1461.
 Holm, R. & Kirschstein, B. 1936 *Wiss. Veröff. Siemenskonz.* **15** (1), 122.
 McFarlane, J. S. & Tabor, D. 1948 *Proc. Seventh Int. Congr. Appl. Mech., Lond.*, **4**, 31.
 McFarlane, J. S. & Tabor, D. 1950 *Proc. Roy. Soc. A*, **202**, 224.
 Moore, A. J. W. 1948 *Proc. Roy. Soc. A*, **195**, 231.
 Nadai, A. 1931 *Plasticity*. New York.
 Whitehead, J. R. 1950 *Proc. Roy. Soc. A*, **201**, 109.

Homotopy invariants and continuous mappings

BY SU-CHENG CHANG

British Council Scholar, Magdalen College, University of Oxford

(Communicated by J. H. C. Whitehead, F.R.S.—Received 11 February 1950)

This paper contributes new numerical invariants to the topology of a certain class of polyhedra. These invariants, together with the Betti numbers and coefficients of torsion, characterize the homotopy type of one of these polyhedra. They are also applied to the classification of continuous mappings of an $(n+2)$ -dimensional polyhedron into an $(n+1)$ -sphere ($n > 2$).

INTRODUCTION

According to J. H. C. Whitehead (1948) an arcwise connected polyhedron, K , is called an A_n^2 -polyhedron if $\dim K \leq n+2$ and $\pi_r(K) = 0$, $r = 1, \dots, n-1$. Throughout this paper we assume $n > 2$. He defines an algebraic A_n^2 -cohomology system as follows. Let H^n, H^{n+1}, H^{n+2} be additive groups, each of which has a finite number of generators,† H^n being free. A new group, $H^n(2)$, is defined as the direct sum‡

$$H_2^n + \Delta^*({}_2H^{n+1}),$$

where $\Delta^*({}_2H^{n+1})$ is the image of ${}_2H^{n+1}$ under an isomorphism Δ^* . Here Δ^* is determined modulo an arbitrary homomorphism

$${}_2H^{n+1} \rightarrow H_2^n.$$

In other words, the isomorphism Δ^* is not determined uniquely, but there is a uniquely determined homomorphism

$$\Delta: H^n(2) \rightarrow {}_2H^{n+1},$$

† Only finite, connected complexes are involved.

‡ Let G be an additive Abelian group then $G_2 = G - 2G$ and ${}_2G$ is the subgroup of G which consists of all the elements g , such that $2g = 0$.

such that $\Delta\Delta^* = 1$ and $\Delta^{-1}(0) = H_2^n$. Furthermore, a homomorphism

$$\mu: H^r \rightarrow H^r(2) \quad (r = n, n+2),$$

on to H_2^r , is defined in the usual way, where $H^{n+2}(2) = H_2^{n+2}$. Finally, there is given a homomorphism

$$\gamma^*: H^n(2) \rightarrow H^{n+2}(2),$$

which may be arbitrary. Then $H^n, H^{n+1}, H^{n+2}, H^n(2), H^{n+2}(2)$, together with μ, Δ and γ^* , constitute an A_n^2 -cohomology system.

In an A_n^2 -polyhedron, K , the groups H^r ($r = n, n+1, n+2$) are the r -dimensional integral cohomology groups of K and $H^r(2)$ ($r = n, n+2$), the r -dimensional cohomology groups of K with integers reduced mod 2 as coefficients, $\Delta = (\frac{1}{2})\delta$ and $\dagger \gamma^*$ is the Steenrod square (Steenrod 1947)

$$\gamma^*x = x \cup_{n-2} x \in H^{n+2}(2),$$

where $x \in H^n(2)$. Therefore a definite cohomology system is associated with K . A homomorphism, f , between two cohomology systems contains five homomorphisms of the corresponding groups. If f commutes with μ and Δ , it is called a $[\mu, \Delta]$ -homomorphism. If f also commutes with γ^* , it is called a *proper homomorphism*. J. H. C. Whitehead (1948) proved that any A_n^2 -cohomology system is properly isomorphic to the cohomology system of some A_n^2 -polyhedron and that two A_n^2 -polyhedra are of the same homotopy type if, and only if, their cohomology systems are properly isomorphic. In § 2 the homotopy type of an A_n^2 -polyhedron, K , is proved, by means of a technique developed in § 1, to be determined by a set of numerical invariants of K , among which are certain new ones, called *secondary torsions*. The secondary torsions, like the Betti numbers and coefficients of torsion, can, in principle, be calculated constructively. In § 3 it is shown that the homotopy group, $\pi_{n+1}(K)$, may be calculated constructively in terms of the Betti numbers, torsions and secondary torsions of K . It follows from theorem 6 in J. H. C. Whitehead (1949b) that the secondary torsions are invariants of the $(n+2)$ -type, as defined in the same paper, and *a fortiori* they are also invariants of the homotopy type of any complex K , such that $\pi_r(K) = 0$, $r = 1, \dots, n-1$ ($n > 2$).

In the last section we compute the cohomotopy group $\pi^{n+1}(K^{n+2})$, if $n > 2$. It will be clear that this group is determined by our numerical invariants.

1. INDUCED AUTOMORPHISMS

1.1. Suppose two cohomology systems

$$H = H(H^n, H^{n+1}, H^{n+2}, H^n(2), H^{n+2}(2), \mu, \Delta, \gamma^*)$$

and

$$\bar{H} = \bar{H}(\bar{H}^n, \bar{H}^{n+1}, \bar{H}^{n+2}, \bar{H}^n(2), \bar{H}^{n+2}(2), \bar{\mu}, \bar{\Delta}, \bar{\gamma}^*)$$

are given. Let $h_r: H^r \approx H^r$, for $r = n, n+1, n+2$. Then the family h_r may be extended $\dagger$ to a (μ, Δ) -isomorphism, $h: H \approx \bar{H}$, and h^{-1} carries $\bar{H}$ into

$$H' = H'(H^n, H^{n+1}, H^{n+2}, H^n(2), H^{n+2}(2), \mu, \Delta, h^{-1}\bar{\gamma}^*h),$$

$\dagger$ The operator $(\frac{1}{2})\delta$ seems to have been used first by Whitney (1937b).

$\ddagger$ This follows from a simplified version of Lemmas 2, 3 in J. H. C. Whitehead (1949a).

where h operates on $H^n(2)$ and h^{-1} on $\bar{H}^{n+2}(2)$. If there is a proper isomorphism $f: H \approx \bar{H}$. Then $f\gamma^* = \bar{\gamma}^*f$ and

$$(h^{-1}f)\gamma^* = h^{-1}\bar{\gamma}^*f = h^{-1}\bar{\gamma}^*h(h^{-1}f),$$

which means that the (μ, Δ) -automorphism $h^{-1}f: H \approx H'$ is a proper isomorphism. Conversely, if $f': H \approx H'$ is proper, so is hf' , since

$$(hf')\gamma^* = hh^{-1}\bar{\gamma}^*hf' = \bar{\gamma}^*(hf').$$

Therefore our algebraic problem is to find necessary and sufficient conditions in order that

$$\gamma^* = f\gamma_1^*f^{-1}: H^n(2) \rightarrow H^{n+2}(2),$$

where $\gamma^*, \gamma_1^*: H^n(2) \rightarrow H^{n+2}(2)$ are given homomorphisms and f is a $[\mu, \Delta]$ -automorphism. In other words, it is a problem of classifying the homomorphisms γ^* under the group of $[\mu, \Delta]$ -automorphisms of H .

Let a cohomology system be given. Then an automorphism, b , of $H^{n+2}(2)$ is called a μ -automorphism if, and only if, $b = \mu g_{n+2} \mu^{-1}$, where g_{n+2} is an automorphism of H^{n+2} . It is single-valued because $\mu g_{n+2} \mu^{-1}(0) = 0$. On the other hand, an automorphism, a , of $H^n(2)$ is called a Δ -automorphism if, and only if, there exist automorphisms g_n, g_{n+1} of H^n, H^{n+1} , such that

$$a\mu = \mu g_n, \quad \Delta a = g'_{n+1} \Delta,$$

where $g'_{n+1} = g_{n+1} |_{2H^{n+1}}$. Obviously $g_n, g_{n+1}, g_{n+2}, a, b$ constitute a $[\mu, \Delta]$ -automorphism g , of H . Conversely, if f is a $[\mu, \Delta]$ -automorphism, then $f|_{H^n(2)}$ is a Δ -automorphism and $f|_{H^{n+2}(2)}$ is a μ -automorphism.

The groups $H^n(2)$ and $H^{n+2}(2)$ are free mod 2 modules. Let bases for these modules be given. Then the automorphisms a, b are given by matrices $[a], [b]$, whose elements are integers, reduced mod 2. When $[a], [b]$ are thus determined we shall call $[a]$ a Δ -matrix and $[b]$ a μ -matrix. Let $[\gamma^*], [\gamma'^*]$ be the (mod 2) matrices by means of which γ^* and $\gamma'^* = b^{-1}\gamma^*a$ are expressed in terms of the given bases for $H^n(2), H^{n+2}(2)$. Then the $[\mu, \Delta]$ -automorphism which corresponds to a, b is represented by the transformation

$$[\gamma^*] \rightarrow [\gamma'^*] = [a][\gamma^*][b]^{-1}. \quad (1.1)$$

Therefore the problem is reduced to the arithmetical one of classifying mod 2 matrices $[\gamma^*]$ under transformations of the form (1.1), where the matrices $[a], [b]$ are subject to certain restrictions described in §§ 1.3, 1.4 below.

1.2. Let $e_1, \dots, e_n$ be generators of n cyclic groups G_l ($l = 1, \dots, n$) of orders $p^{r_1}, \dots, p^{r_n}$ respectively, where p is a prime number and $r_1 \leq r_2 \leq \dots \leq r_n$. Let $G = \Sigma G_l$ be the direct sum of all the groups, G_l . An endomorphism $E: G \rightarrow G$ is given by

$$Ee_l = \sum_{i=1}^n \alpha_{li} e_i, \quad (1.2)$$

where α_{li} are integers satisfying

$$p^{r_l} \alpha_{li} \equiv 0, \quad \text{mod } p^{r_i}. \quad (1.3)$$

LEMMA 1. E is an automorphism of G if, and only if, $|\alpha_i|$ is prime to p .

Proof. Let $A = |\alpha_i|$, $(A, p) = 1$. Take β_{li} to be the algebraic complement of α_{li} in $|\alpha_{li}|$. Then

$$\beta_{li} = (-1)^{l+i} \begin{vmatrix} \alpha_{11} & \dots & \alpha_{1l-1} & \alpha_{1l+1} & \dots & \alpha_{1n} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \alpha_{i-11} & \dots & \alpha_{i-1l-1} & \alpha_{i-1l+1} & \dots & \alpha_{i-1n} \\ \alpha_{i+11} & \dots & \alpha_{i+1l-1} & \alpha_{i+1l+1} & \dots & \alpha_{i+1n} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \alpha_{n1} & \dots & \alpha_{nl-1} & \alpha_{nl+1} & \dots & \alpha_{nn} \end{vmatrix}.$$

If $l = i$, the condition (1.3) is vacuous. If $l \neq i$, β_{li} may be written as a sum of terms of the following form:

$$\dots \alpha_{ln_s} \alpha_{n_s n_s - 1} \dots \alpha_{n_1 n_0} \alpha_{n_0 l} \dots$$

Repeated use of (1.3) shows that

$$p^{r_i} \beta_{li} \equiv 0, \quad \text{mod } p^{r_i}.$$

Therefore an endomorphism of G is defined by

$$E^0 e_l = \sum_{i=1}^n \beta_{li} e_i \quad (l = 1, \dots, n).$$

Since $[\alpha_{li}] [\beta_{kl}] = AI$, it follows that

$$E^0 E e_l = E E^0 e_l = A e_l.$$

Since $(A, p) = 1$, it follows that $E^0 E: G \approx G$, $E E^0: G \approx G$. Therefore $E: G \approx G$.

Conversely, if E^{-1} is given by $E^{-1} e_l = \sum_i \beta'_{li} e_i$,

where $p^{r_i} \beta'_{li} \equiv 0 \pmod{p^{r_i}}$, then

$$A |\beta'_{li}| \equiv 1, \quad \text{mod } p.$$

Therefore $(A, p) = 1$ and the lemma is proved.

Let I_0 be the group of integers and I_m , the (additive) group of residue classes mod m . Let $\mu: I_0 \rightarrow I_2$ be the natural homomorphism.

LEMMA 2. If $[a_{ij}]$ is a square matrix of integers whose determinant is odd, there is a unimodular matrix of integers, $[a'_{ij}]$, such that $a_{ij} - a'_{ij} \equiv 0 \pmod{2}$, for all i and j .

Proof. Since $|a_{ij}|$ is odd, the matrix $[\mu a_{ij}]$ is a unimodular mod 2 matrix, which can be reduced to the unit mod 2 matrix by a series of elementary transformations, each of which consists of adding one column to another. Therefore $[\mu a_{ij}]$ is a product of matrices of the form $[\mu(I + e_{pq})]$, where I is the unit matrix and the matrix e_{pq} has all its elements zero except for a 1 in the p th row and q th column ($p \neq q$). Evidently $\mu a_{ij} = \mu a'_{ij}$, where $[a'_{ij}]$ is the corresponding product of the (unimodular) matrices $[I + e_{pq}]$. This proves the lemma.

1.3. All our Abelian groups will be represented as direct sums of cyclic groups, whose orders are either infinite or powers of primes. Let $f_1, \dots, f_{\mu_0}$ be the free generators of H^{n+2} and let

$$e(1), \dots, e(S_1); \quad e(S_1 + 1), \dots, e(S_2); \quad \dots; \quad \dots e(S_N)$$

be those generators of H^{n+2} whose (finite) orders are even. More precisely, let $e(S_{l-1}+1), \dots, e(S_l)$ be of order 2^{m_l} , where $S_0 = 0$ and $m_{l-1} > m_l$. Now $H^{n+2}(2)$ is a free mod 2 module with generators $\bar{f}_k$ ($k = 1, \dots, \mu_0$) and $\bar{e}(S_{l-1}+u_{l-1})$ ($l = 1, \dots, N$; $u_{l-1} = 1, \dots, \mu_l = S_l - S_{l-1}$) such that

$$\bar{f}_k = \mu f_k, \quad \bar{e}(k) = \mu e(k). \quad (1.4)$$

We consider automorphisms of H^{n+2} . Those generators, which have odd order, do not interest us, and we shall not mention them explicitly even if they are concerned. Let B be a given automorphism of H^{n+2} . Then

$$\left. \begin{aligned} B(f_i) &= \sum_{j=1}^{\mu_0} \alpha_{ij} f_j + \sum_{k=1}^{S_N} \beta_{ik} e(k) + \dots, \quad (i = 1, \dots, \mu_0), \\ B(e(l)) &= \sum_{k=1}^{S_N} \gamma_{lk} e(k) \quad (l = 1, \dots, S_N), \end{aligned} \right\} \quad (1.5)$$

where $[\alpha_{ij}]$ is unimodular, $[\beta_{ik}]$ is arbitrary, the determinant $|\gamma_{lk}|$ is odd and, writing the order of $e(k)$ as σ_k ,

$$\sigma_l \gamma_{lk} \equiv 0, \quad \text{mod } \sigma_k. \quad (1.6)$$

$$\text{This implies} \quad \gamma_{lk} \equiv 0, \quad \text{mod } 2 \quad \text{if } \sigma_k > \sigma_l. \quad (1.7)$$

The automorphism (1.5) induces a μ -automorphism in $H^{n+2}(2)$, given by

$$\left. \begin{aligned} b(\bar{f}_i) &= \sum_{j=1}^{\mu_0} \alpha_{ij} \bar{f}_j + \sum_{k=1}^{S_N} \beta_{ik} \bar{e}(k) \quad (i = 1, \dots, \mu_0), \\ b(\bar{e}(l)) &= \sum_{k=1}^{S_N} \gamma_{lk} \bar{e}(k) \quad (l = 1, \dots, S_N), \end{aligned} \right\} \quad (1.8)$$

subject to (1.7). Conversely, let (1.8) be an automorphism of $H^{n+2}(2)$ which satisfies (1.7). Then $|\alpha_{ij}| \cdot |\gamma_{jk}|$, and hence $|\alpha_{ij}|$, is odd. Lemma 2 gives us a unimodular matrix $[\alpha'_{ij}]$ of integers such that

$$\alpha'_{ij} - \alpha_{ij} \equiv 0, \quad \text{mod } 2.$$

Let $\gamma'_{lk} = \gamma_{lk}$ or $(\sigma_k/\sigma_l)\gamma_{lk}$ according as $\sigma_k \leq \sigma_l$ or $\sigma_k > \sigma_l$. Then the automorphism (1.8) is identical with

$$\left. \begin{aligned} b(\bar{f}_i) &= \sum_{j=1}^{\mu_0} \alpha'_{ij} \bar{f}_j + \sum_{k=1}^{S_N} \beta_{ik} \bar{e}(k) \quad (i = 1, \dots, \mu_0), \\ b(\bar{e}(l)) &= \sum_{k=1}^{S_N} \gamma'_{lk} \bar{e}(k) \quad (l = 1, \dots, S_N), \end{aligned} \right\} \quad (1.9)$$

because γ_{lk} is even if $\sigma_k > \sigma_l$. Evidently (1.9) can be realized by an automorphism of H^{n+2} . Therefore (1.8) or (1.9) is a μ -automorphism.

All the μ -automorphisms of $H^{n+2}(2)$ constitute a group, which will be called the μ -group. We have proved

LEMMA 3. *The μ -group consists of all automorphisms given by (1.8) subject to (1.7).*

The form of a μ -matrix with coefficients reduced mod 2, is indicated by figure 1:

	$\bar{f}_1, \dots, \bar{f}_{\mu_0}$	$\bar{e}(1), \dots, \bar{e}(\mu_1), \dots$	$\bar{e}(S_N)$
$\bar{f}_1$	α	β	β
$\vdots$			
$\bar{f}_{\mu_0}$			
$\bar{e}(1)$	0	γ	γ
$\vdots$			
$\bar{e}(\mu_1)$			
$\vdots$	0	0	γ
$\bar{e}(S_N)$	0	0	γ

FIGURE 1

1.4. Now we deal with $H^n(2)$ in a similar way. Let $(d_1, \dots, d_{\Delta_0})$ be a set of free generators of H^n and let $c(1), \dots, c(r_1); c(r_1 + 1), \dots, c(r_2); \dots, c(r_D)$ be those generators of H^{n+1} , whose orders are powers of 2. Let $c(r_{i-1} + 1), \dots, c(r_i)$ be of order 2^{n_i} , where $n_{i-1} > n_i$, $r_0 = 0 < i \leq D$. We write $r_1 = \Delta_1$, $r_j - r_{j-1} = \Delta_j$ ($j = 2, \dots, D$). Then ${}_2H^{n+1}$ is generated by $\rho_p c(p)$, where $2\rho_p$ is the order of $c(p)$. Let $H_0^{n+1} \subset H^{n+1}$ be the subgroup generated by $c(1), \dots, c(r_D)$. Then $g_{n+1} H_0^{n+1} = H_0^{n+1}$ if $g_{n+1}: H^{n+1} \approx H^{n+1}$. Conversely any automorphism of H_0^{n+1} can be extended to an automorphism of H^{n+1} , since H_0^{n+1} is a direct summand of H^{n+1} . Let $g_n: H^n \approx H^n$, $g_{n+1}: H_0^{n+1} \approx H_0^{n+1}$ be automorphisms given by

$$g_n(d_i) = \sum_{j=1}^{\Delta_0} \epsilon_{ij} d_j, \quad g_{n+1}(c(p)) = \sum_{q=1}^{r_D} \eta_{pq} c(q), \quad (1.10)$$

where $|\epsilon_{ij}| = 1$, $|\eta_{pq}|$ is odd and $2\rho_p \eta_{pq} \equiv 0 \pmod{2\rho_q}$, whence $\rho_p \eta_{pq} \equiv 0 \pmod{\rho_q}$. Let Δ^* be any right inverse of Δ and let $a: H^n(2) \approx H^n(2)$ be given by

$$\left. \begin{aligned} a\mu d_i &= \sum_{j=1}^{\Delta_0} \epsilon_{ij} \mu d_j, \\ a\Delta^* \rho_p c(p) &= \sum_{j=1}^{\Delta_0} \zeta_{pj} \mu d_j + \sum_{q=1}^{r_D} \bar{\eta}_{pq} \Delta^* \rho_q c(q), \end{aligned} \right\} \quad (1.11)$$

where $\bar{\eta}_{pq} = (\rho_p/\rho_q) \eta_{pq}$ and $[\zeta_{pj}]$ is arbitrary. Since $a\mu d_j = \mu g_n d_j$, and

$$\begin{aligned} \Delta a \Delta^* \rho_p c(p) &= \sum_{q=1}^{r_D} \bar{\eta}_{pq} \Delta \Delta^* \rho_q c(q) = \sum_{q=1}^{r_D} \bar{\eta}_{pq} \rho_q c(q) \\ &= \sum_{q=1}^{r_D} \rho_p \eta_{pq} c(q) = \rho_p g_{n+1} c(p) \\ &= g_{n+1} \rho_p c(p) = g_{n+1} \Delta \Delta^* \rho_p c(p), \end{aligned}$$

it follows that a is a Δ -automorphism. Since $\bar{\eta}_{pq} = (\rho_p/\rho_q) \eta_{pq}$ it follows that

$$\bar{\eta}_{pq} \equiv 0, \pmod{2} \quad \text{if} \quad \rho_p > \rho_q. \quad (1.12)$$

Conversely, let (1.11) be a given automorphism of $H^n(2)$, which satisfies (1.12). For similar reasons to those in the proof of Lemma 1 it may be assumed that $|\epsilon_{ij}| = 1$ and $\bar{\eta}_{pq} = (\rho_p/\rho_q)\eta_{pq}$ if $\rho_p \geq \rho_q$ with integral values of η_{pq} . Let this be so and let $\eta_{pq} = (\rho_q/\rho_p)\bar{\eta}_{pq}$ if $\rho_p < \rho_q$. Then $|\eta_{pq}| = (\rho/\rho)|\bar{\eta}_{pq}| = |\bar{\eta}_{pq}|$, where $\rho = \rho_1 \dots \rho_D$. Therefore $|\eta_{pq}|$ is odd and (1.10), with these values of ϵ_{ij} , η_{pq} , is related as above to (1.11). Therefore (1.11) is a Δ -automorphism.

The group of Δ -automorphisms of $H^n(2)$ will be called the Δ -group. We have proved:

LEMMA 4. *The Δ -group consists of all the automorphisms of the form (1.11), which satisfy (1.12).*

The form of a Δ -matrix, with coefficients reduced mod 2, is indicated by figure 2:

	$\bar{d}_1, \dots, \bar{d}_{\Delta_0}, \bar{c}(1), \dots, \bar{c}(\Delta_1), \dots$		$\bar{c}(r_D)$
$\bar{d}_1$	ϵ	0	0
$\vdots$			
$\bar{d}_{\Delta_0}$			
$\bar{c}(1)$	ζ	$\bar{\eta}$	0
$\vdots$			
$\bar{c}(\Delta_1)$	ζ	$\bar{\eta}$	0
$\bar{c}(r_D)$	ζ	$\bar{\eta}$	$\bar{\eta}$

FIGURE 2

2. THE SECONDARY TORSIONS

If a Δ -automorphism, a^{-1} , and a μ -automorphism, b , are given, the proofs of lemma 3 and lemma 4 imply the existence (not necessarily unique) of a $[\mu, \Delta]$ -automorphisms, f , of the cohomology system such that

$$f|H^n(2) = a^{-1}, \quad f|H^{2+n}(2) = b,$$

and that a^{-1} and b are induced by $f|H^r$ ($r = n, n+1, n+2$). For the reasons given in § 1 our classification problem is equivalent to the classification of (mod 2) matrices $[\gamma^*]$ under the transformation (1.1), where $[a]$, $[b]$ are (mod 2) matrices of the form indicated in figures 2 and 1. Let $\kappa(l)$ ($l = 0, 1, \dots, N$), $\rho(i)$ ($i = 0, 1, \dots, D$) be the submatrices of $[\gamma^*]$, which consist of the columns numbered $\mu_0 + s_{l-1} + 1, \dots, s_l + \mu_0$ ($s_{l-1} = -\mu_0$, $s_l = 0$ if $l = 0$) and of the rows numbered $\Delta_0 + r_{i-1} + 1, \dots, \Delta_0 + r_i$ ($r_{i-1} = -\Delta_0$, $r_i = 0$ if $i = 0$). Thus $\rho(i)$, $\kappa(l)$ correspond to the $(i+1)$ th column of blocks in figure 2 and the $(l+1)$ th row of blocks in figure 1. Let $\gamma^*(i, l)$ be the submatrix of $[\gamma^*]$ in which $\rho(i)$ and $\kappa(l)$ intersect, and let Γ_{il} be the submatrix in which the first $\Delta_0 + r_i$ rows and the first $\mu_0 + s_l$ columns intersect. Thus Γ_{il} consists of all the blocks $\gamma^*(p, q)$ for $p \leq i$, $q \leq l$. It is evident from figures 1 and 2 that (1.1) is the resultant of elementary transformations, each of which consists of adding (mod 2) a column in $\kappa(l)$ to a column in $\kappa(l_1)$, where $l_1 \geq l$, or a row in $\rho(i)$ to a row in $\rho(i_1)$,

where $i_1 \geq i$. These elementary transformations do not alter the rank, T_{il} , of Γ_{il} . Therefore T_{il} , and likewise

$$\tau_{il} = T_{il} - T_{il-1} - T_{i-1l} + T_{i-1l-1}$$

($T_{il} = 0$ if $i < 0$ or $l < 0$), are invariant under transformations of the form (1.1). We call the numbers τ_{il} the *secondary torsions* of the polyhedron or cohomology system to which γ^* belongs.

The matrix $[\gamma^*]$ can obviously be reduced, by the elementary transformations described in the preceding paragraph, to one in which there is at most one unity in each row and in each column. After this reduction the number τ_{il} will be the rank of the block $\gamma^*(i, l)$. Finally $[\gamma^*]$ can be reduced to a unique normal form, by permuting rows and columns within each strip $\rho(i)$ or $\kappa(l)$, in which the units in $\gamma^*(i, l)$, if any, occur in the first τ_{il} rows and columns of $\gamma^*(i, l)$.

Two A_n^2 -cohomology systems are properly isomorphic if, and only if, the corresponding groups H^r have the same rank and invariant factors, for each $r = n, n+1, n+2$, and if the secondary torsions are the same in both. For in this case $H^r, H^n(2), H^{n+2}(2)$ may be referred to bases such that the homomorphisms μ, Δ, γ^* are represented by the same matrices in both systems. Hence we have the

THEOREM 1. *The homotopy type of an A_n^2 -polyhedron ($n > 2$) is determined by its Betti numbers, torsion coefficients and secondary torsions.*

If H^{n+1} and H^{n+2} are free then we have m different homotopy types where $m = \min(\Delta_0, \mu_0)$. If $\pi_{n+1}(K) = 0$, then γ^* is an isomorphism into (see J. H. C. Whitehead 1948), all of the secondary torsions reach their maximum values and the homotopy type is unique.

3. NORMAL A_n^2 -POLYHEDRA

Let a cohomology system be given, with a normalized γ^* -matrix. We shall construct a normal polyhedron whose cohomology system is properly isomorphic to the given one. By an elementary A_n^2 -polyhedron we mean one of the following kinds:

$$B_1^r = S^r \quad (r = n, n+1, n+2).$$

$B_2(\sigma) = S^n \cup e^{n+1}$, where e^{n+1} is attached to S^n by a map $f: \dot{E}^{n+1} \rightarrow S^n$ of degree σ , a power of a prime.

$B_3(\tau) = S^{n+1} \cup e^{n+2}$, where e^{n+2} is attached to S^{n+1} by a map $f: \dot{E}^{n+2} \rightarrow S^{n+1}$ of degree τ , a power of a prime.

$$B_4 = S^n \cup e^{n+2}, \text{ where } e^{n+2} \text{ is attached to } S^n \text{ by an essential map } f: \dot{E}^{n+2} \rightarrow S^n.$$

$B_5(2^p) = S^n \cup e^{n+1} \cup e^{n+2}$ where e^{n+1} is attached to S^n by a map $f: \dot{E}^{n+1} \rightarrow S^n$ of degree 2^p and e^{n+2} is attached to S^n by an essential map $f: \dot{E}^{n+2} \rightarrow S^n$.

$B_6(2^q) = S^n \cup S^{n+1} \cup e^{n+2}$, where $S^n \cap S^{n+1} = p_0$, a point, and e^{n+2} is attached to $S^n \cup S^{n+1}$ by a map, $f: \dot{E}^{n+2} \rightarrow S^n \cup S^{n+1}$, of the form $a + b$, where a denotes an essential map of $\dot{E}^{n+2}$ on to S^n and b maps $\dot{E}^{n+2}$ on to S^{n+1} with degree 2^q .

$$B_7(2^p, 2^q) = B_6(2^q) \cup e^{n+1}. \text{ Attach } e^{n+1} \text{ to } B_6(2^q) \text{ by a map } f: \dot{E}^{n+1} \rightarrow S^n \text{ of degree } 2^p.$$

In the first three cases it is unnecessary to consider the normal γ^* -matrices, because either $H^n(2)$ or $H^{n+2}(2)$ is trivial. It follows from theorem 4 in J. H. C. Whitehead (1948) that the normal γ^* -matrix in any other case is the unit matrix with a single row and column.

We now take arbitrary numbers†, $Y_{1r}, Y_2(\sigma), \dots, Y_7(2^p, 2^q)$, of polyhedra of types $B_1^r, B_2(\sigma), \dots, B_7(2^p, 2^q)$. To begin with these polyhedra shall be disjoint. We choose a point on the sphere of lowest dimensionality in each of them and identify all these points with a new point p_0 , thus forming an A_n^2 -polyhedron K . Obviously $Y_{1n+1}, Y_2(\sigma), Y_3(\tau)$, with odd values of σ, τ , are the $(n+1)$ st Betti numbers and the odd coefficients of n -dimensional and $(n+1)$ -dimensional torsion. Let $\mu_l, \Delta_i, 2^{m_l}, 2^{n_i}$ mean the same as in § 1.2, with reference to the cohomology system of K . Then

$$\begin{aligned}\Delta_0 &= Y_{1n} + Y_4 + \sum_{l=1}^N Y_6(2^{m_l}), \\ \Delta_i &= Y_2(2^{n_i}) + Y_5(2^{n_i}) + \sum_{l=1}^N Y_7(2^{n_i}, 2^{m_l}), \\ \mu_0 &= Y_{1n+2} + Y_4 + \sum_{i=1}^D Y_5(2^{n_i}), \\ \mu_l &= Y_3(2^{m_l}) + Y_6(2^{m_l}) + \sum_{i=1}^D Y_7(2^{n_i}, 2^{m_l}).\end{aligned}\tag{3.1}$$

Also it follows from theorem 4 in J. H. C. Whitehead (1948) that

$$\begin{aligned}\tau_{00} &= Y_4, & \tau_{0l} &= Y_6(2^{m_l}), \\ \tau_{i0} &= Y_5(2^{n_i}), & \tau_{il} &= Y_7(2^{n_i}, 2^{m_l}).\end{aligned}\tag{3.2}$$

Therefore the equations (3.1) may be written as

$$\begin{aligned}Y_2(2^{n_i}) &= \Delta_i - \sum_{\beta=0}^N \tau_{i\beta} \quad (i = 0, \dots, D), \\ Y_3(2^{m_l}) &= \mu_l - \sum_{\alpha=0}^D \tau_{\alpha l} \quad (l = 0, \dots, N),\end{aligned}\tag{3.3}$$

where $Y_2(2^{n_0}) = Y_{1n}$, $Y_3(2^{m_0}) = Y_{1n+2}$. Thus the equations (3.1) and (3.2) can be solved uniquely for the Y 's. Since τ_{il} is the rank of $\gamma^*(i, l)$ when γ^* has its normal form, and since there is then at most one unity in each row or column of γ^* , it follows that $Y_2(2^{n_i}) \geq 0$, $Y_3(2^{m_l}) \geq 0$. Hence it follows that K , with the Y 's suitably chosen, is a normal form for the homotopy type of a given A_n^2 -polyhedron. If we lay down a definite rule for attaching the cells to the spheres in $B_2(\sigma)$ etc., the complex K will be unique up to a (combinatorial) isomorphism. We shall describe K as a normal A_n^2 -complex.

4. THE GROUP π_{n+1}

Let P be a given A_n^2 -polyhedron and let K be a normal A_n^2 -polyhedron of the same homotopy type as P . Then $\pi_{n+1}(P) \approx \pi_{n+1}(K)$. It follows from Hilton (1950) that $\pi_{n+1}(K)$ is the direct sum of the groups $\pi_{n+1}(K_i)$, where $K_1, K_2, \dots$ are the elementary polyhedra of which K is composed.

THEOREM 2. *The groups $\pi_{n+1}(K_i)$ are given by table 1, in which Z_q denotes a cyclic group of order q ($q \leq \infty$) and the groups not shown are all zero.*

† Some or all of the Y 's may be zero. If all are zero, then K (below) is a single point.

TABLE 1

	B_1^n	B_1^{n+1}	$B_2(2r)$	$B_3(\tau)$	$B_6(2^p)$	$B_7(2^p, 2^q)$
π_{n+1}	Z_2	Z_∞	Z_2	Z_7	$Z_{2^{p+1}}$	$Z_{2^{p+1}}$

This is an obvious consequence of Hilton (1950). Therefore† $\pi_{n+1}(K)$ and hence $\pi_{n+1}(P)$ is given by this table and (3.2), (3.3) when the numerical invariants of K have been calculated.

5. THE COHOMOTOPY GROUP $\pi^{n+1}(K)$

Let P be an arbitrary (finite) complex and let P_0 be the complex which is formed by shrinking P^{n-1} to a point ($n > 2$). The homotopy classes of maps $f: P \rightarrow S^{n+1}$ are obviously in (1-1)-correspondences with the classes of maps $g: P_0 \rightarrow S^{n+1}$ in such a way that $\pi^{n+1}(P) \approx \pi^{n+1}(P_0)$. Moreover, $\pi^{n+1}(P_0) \approx \pi^{n+1}(K)$, where K is a normal A_n^2 -complex of the same homotopy type as P_0 . Obviously $\pi^{n+1}(K)$ is the direct sum of the groups $\pi^{n+1}(K_i)$, where $K_1, K_2, \dots$ are the elementary polyhedra of which K is composed.

THEOREM 3. *The groups $\pi^{n+1}(K_i)$ are given by table 2, those not shown being zero.*

TABLE 2

	B_1^{n+1}	B_1^{n+2}	$B_2(\sigma)$	$B_3(2r)$	$B_5(2^p)$	$B_7(2^p, 2^q)$
π^{n+1}	Z_∞	Z_2	Z_σ	Z_2	$Z_{2^{p+1}}$	$Z_{2^{p+1}}$

The number of elements in each group $\pi^{n+1}(K_i)$ may be computed without difficulty by means of theorems‡ in Hopf (1933) and Steenrod (1947). It remains to prove that all these groups are cyclic (or zero).

As observed in Spanier (1949) it is obvious that $\pi^{n+1}(B_1^r) = 0, Z_\infty$, or Z_2 according as $r = n, n+1$ or $n+2$. Let K_i be of any of the other types, $B_2(\sigma), \dots$, and let $L_i \subset K_i$ be the subcomplex which is indicated in table 3, with $L_i = K_i$ if $K_i = B_4$ or $B_6(2^q)$. Let $f: K_i \rightarrow S^{n+1}$ be a given map. It follows from theorems in Hopf (1933) and Steenrod (1947) that $f|_{L_i}$ is homotopic to a constant map. Therefore the homomorphism, $\pi^{n+1}(K_i, L_i) \rightarrow \pi^{n+1}(K_i)$, which is induced by the identical map, $K_i \rightarrow (K_i, L_i)$, is on to $\pi^{n+1}(K_i)$. Since $\pi^r(S^{n+1})$ is cyclic§ if $r \leq 2n+2$ it is easily seen that $\pi^{n+1}(K_i, L_i)$, and hence $\pi^{n+1}(K_i)$ is cyclic (or zero). This proves the theorem.

TABLE 3

K_i	$B_2(\sigma)$	$B_3(\tau)$	$B_5(2^p)$	$B_7(2^p, 2^q)$
L_i	S^n	S^{n+1}	$S^n \cup e^{n+2}$	$S^n \cup S^{n+1} \cup e^{n+2}$

I owe a great deal to Professor J. H. C. Whitehead and wish to thank him for his inspiring suggestions and valuable criticisms.

† The algebraic construction of $\pi_n(K)$ is given by a theorem due to Hurewicz (1936), which states that $\pi_n(K) \approx H_n(K)$. A relation between $\pi_{n+1}(K)$ and $H_{n+1}(K)$ has appeared in G. W. Whitehead (1948).

‡ Also see Eilenberg (1940) and Whitney (1937a).

§ See Spanier (1949) and Hu (1949).

REFERENCES

- Eilenberg, S. 1940 *Ann. Math. Princeton, etc.* **41**, 231–251.
 Freudenthal, H. 1937 *Compositio Math.* **5**, 299–314.
 Hilton, P. J. 1950 *Quart. J. Math.* (in the press).
 Hopf, H. 1933 *Comment. math. helv.* **5**, 39–54.
 Hopf, H. 1945 *Comment. math. helv.* **17**, 307–326.
 Hu, S. T. 1949 *Ann. Math. Princeton, etc.* **50**, 158–173.
 Hurewicz, W. 1936 *Proc. Akad. Sci. Amst.* **39**, 117–125.
 Spanier, E. 1949 *Ann. Math. Princeton, etc.* **50**, 203–245.
 Steenrod, N. E. 1947 *Ann. Math. Princeton, etc.* **48**, 290–320.
 Whitehead, G. W. 1948 *Proc. Nat. Acad. Sci., Wash.*, **34**, 207–211.
 Whitehead, J. H. C. 1948 *Ann. Soc. Math. Polonaise*, **21**, 176–186.
 Whitehead, J. H. C. 1949 *Comment. math. helv.* **22**, 48–92.
 Whitehead, J. H. C. 1949b *Bull. Amer. Math. Soc.* **55**, 213–245.
 Whitney, H. 1937a *Duke Math. J.* **3**, 51–55.
 Whitney, H. 1937b *Bull. Amer. Math. Soc.* **43**, 785–805.

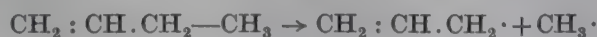
The $\text{CH}_2:\text{CH}.\text{CH}_2-\text{CH}_3$ bond dissociation energy and the heat of formation of the allyl radical

BY A. H. SEHON AND M. SZWARC

Chemistry Department, University of Manchester

(Communicated by M. G. Evans F.R.S.—Received 16 February 1950)

The pyrolysis of butene-1 was investigated by a flow technique, toluene being used as a carrier gas. It was found that butene-1 decomposed into allyl and methyl radicals according to the equation



Methyl radicals were removed by reaction with toluene giving methane and benzyl radicals. The rate of the initial decomposition was measured by the rate of formation of methane.

The decomposition was found to be a homogeneous first order gas reaction. The activation energy was calculated at 61.5 kcal./mole and it was identified with the $\text{CH}_2:\text{CH}.\text{CH}_2-\text{CH}_3$ bond dissociation energy. Taking $D(\text{CH}_2:\text{CH}.\text{CH}_2-\text{CH}_3)$ at 61.5 kcal./mole we calculated from thermochemical data $D(\text{CH}_2:\text{CH}.\text{CH}_2-\text{H})$ at 76.5 kcal./mole and the heat of formation of allyl radical at +30 kcal./mole.

The fate of allyl radicals is discussed and the thermal stability of these is compared with that of benzyl radicals.

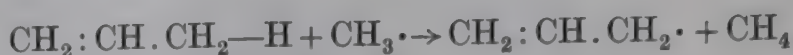
INTRODUCTION

For a number of years it has been recognized that the allyl radical is stabilized by its high resonance energy. This is substantiated both on theoretical and experimental grounds.

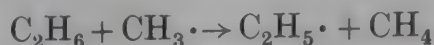
As a result of various quantum-mechanical calculations the resonance energy of the allyl radical was estimated as 15.4 kcal./mole (Coulson 1938), and as 18.7 kcal./mole (Orr 1946).

The effect of resonance energy in the radical should manifest itself by a decrease of the C—X bond dissociation energy in any $\text{CH}_2:\text{CH}.\text{CH}_2\text{—X}$ molecule as compared with the respective bond dissociation energy in $\text{CH}_3\text{—X}$. In agreement with this expectation Butler & Polanyi (1943) reported that the rate of decomposition of allyl iodide was 4.5×10^4 times faster than that of methyl iodide at 430°C . Similarly, it was observed in these laboratories that the decomposition of allyl bromide is 3×10^5 times faster than that of methyl bromide at 500°C .

The stability of the allyl radical and the relative weakness of the $\text{CH}_2:\text{CH}.\text{CH}_2\text{—H}$ bond is indicated also by the ability of propylene to inhibit chain reactions (Rice & Polly 1938, Smith & Hinshelwood 1942). The same conclusion follows from the low activation energy of 3.1 kcal./mole required for the reaction between methyl radicals and propylene



in comparison with the reactions

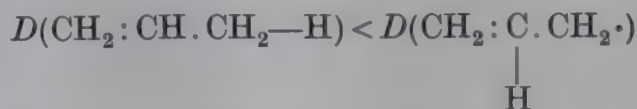


which require 8 kcal./mole (Taylor & Smith 1940).

On the basis of all this evidence it was reasonable to assume that the weakest bond in propylene is the $\text{CH}_2:\text{CH}.\text{CH}_2\text{—H}$ bond. In order to confirm this assumption Szwarc (1949a) investigated the thermal decomposition of propylene. It was hoped that propylene would dissociate into hydrogen and allyl radicals, and that the latter on account of their stability would be removed as such from the reaction zone. This unimolecular decomposition process would have then led to the determination of the C—H bond dissociation energy in propylene and to the direct evaluation of the heat of formation of the allyl radical from the equation

$$\Delta H_f(\text{allyl}) = \Delta H_f(\text{CH}_2:\text{CH}.\text{CH}_3) + D(\text{CH}_2:\text{CH}.\text{CH}_2\text{—H}) - \Delta H_f(\text{H}).$$

It must be emphasized that the essential condition for allyl radicals not to undergo further decomposition in the reaction zone is



From the prevailing thermochemical data $\Delta H_f(\text{CH}_2:\text{CH}.\text{CH}_3) = 4.9$ kcal./mole and $\Delta H_f(\text{CH}_2:\text{C}:\text{CH}_2) = 45.9$ kcal./mole, we obtain the following relationship

$$D(\text{CH}_2:\text{CH}.\text{CH}_2\text{—H}) + D(\text{CH}_2:\underset{\text{H}}{\underset{|}{\text{C}}}.\text{CH}_2\cdot) = 145 \text{ kcal./mole.}$$

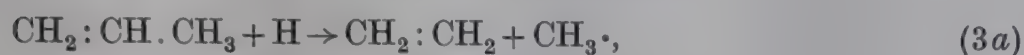
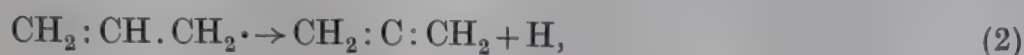
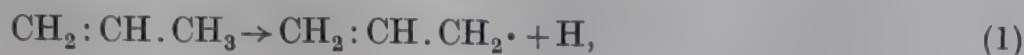
It is obvious, therefore, that the above condition will be satisfied if

$$D(\text{CH}_2:\text{CH}.\text{CH}_2\text{—H}) < 72 \text{ kcal./mole.}$$

It seems, however, that this value is much too low. Indeed, the kinetics of the thermal decomposition of propylene, supplemented by the results obtained from

the pyrolysis of isobutene (Szwarc 1949*b*), suggest that although allyl radicals are formed in the first step of the decomposition of propylene, they dissociate further into allene and hydrogen atoms giving rise to a chain reaction. In spite of this it was possible to deduce from the assumed mechanism of reaction that the $D(\text{C}-\text{H})$ in propylene was about 78 kcal./mole.

It will be sufficient at this stage to summarize briefly the reasoning that led to this value. The overall decomposition of propylene was shown to be a homogeneous gas reaction of the first order, the first order rate constant being given by the expression $1.1 \times 10^{13} \exp(-72000/RT)$. The experimental results could be best interpreted in terms of the following mechanism



This mechanism represents a chain process in which reaction (1) is the chain initiation; reactions (2) and (3), or alternatively (2), (3*a*) and (3*b*), are the chain propagation; and the reaction (4) is the chain termination.

Applying the stationary state method it was shown that the rate of reaction is given by the expression

$$d(\text{H}_2 + \text{CH}_4)/dt = \left(\frac{k_1 k_2 k_3}{k_4} \right)^{\frac{1}{2}} [\text{propylene}]$$

the length of the chain being equal to

$$\left(\frac{k_2 k_3}{k_1 k_4} \right)^{\frac{1}{2}}.$$

Reactions (1) and (2) are unimolecular decompositions of propylene and allyl radicals respectively. It was assumed, therefore, that both should have a frequency factor of the order of $10^{13} \text{ sec.}^{-1}$. Making the usual assumption that the recombination of radicals involves no activation barrier, it follows that the activation energies of reactions (1) and (2) are the above mentioned $\text{CH}_2:\text{CH}.\text{CH}_2-\text{H}$ and $\text{CH}_2:\text{C}.\text{CH}_2\cdot$



bond dissociation energies. On the other hand, considering that the collision frequencies for the bimolecular reactions (3) and (4) have the same value, and that the activation of (3) is very small whereas that of (4) is zero, it was assumed in the first approximation that $k_3 \simeq k_4$. The expression for the rate of reaction is then simplified to

$$\text{rate of reaction} = (k_1 k_2)^{\frac{1}{2}} [\text{propylene}],$$

which leads to the activation energy for the overall process

$$E = \frac{1}{2}[D(\text{CH}_2:\text{CH}.\text{CH}_2-\text{H}) + D(\text{CH}_2:\underset{\text{H}}{\underset{|}{\text{C}}}.\text{CH}_2\cdot)]$$

$$= \frac{145}{2} \text{ kcal./mole} = 72.5 \text{ kcal./mole},$$

and to the frequency factor $(\nu_1 \nu_2)^{\frac{1}{2}} \approx 10^{13} \text{ sec.}^{-1}$. These numerical values are in full agreement with the experimental results.

Similarly, the expression for the length of the chain is reduced to

$$\text{chain length} = (k_2/k_1)^{\frac{1}{2}}.$$

Since on statistical grounds the frequency factor of k_1 is three times as great as that of k_2 it follows that

$$E_1 - E_2 = 2RT[\ln(\text{chain length}) + \frac{1}{2} \ln 3]$$

The actual chain length was assumed to be equal to 10, this assumption being justified on the basis of further results obtained from the pyrolysis of isobutene (Szwarc 1949*b*). It follows, therefore, that

$$E_1 - E_2 = D(\text{CH}_2:\text{CH}.\text{CH}_2-\text{H}) - D(\text{CH}_2:\underset{\text{H}}{\underset{|}{\text{C}}}.\text{CH}_2\cdot)$$

$$= 2RT[\ln 10 + \frac{1}{2} \ln 3] = 11 \text{ kcal./mole},$$

this in conjunction with the previously mentioned relation

$$E_1 + E_2 = D(\text{CH}_2:\text{CH}.\text{CH}_2-\text{H}) + D(\text{CH}_2:\underset{\text{H}}{\underset{|}{\text{C}}}.\text{CH}_2\cdot) = 145 \text{ kcal./mole}$$

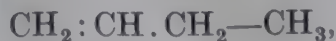
leads to the value of about 78 kcal./mole for the C—H bond dissociation energy in propylene, and to about 67 kcal./mole for the C—H bond dissociation energy in the allyl radical.

In view of the arbitrarily chosen values for the chain length and for the frequency factors of the individual steps, it was thought desirable to obtain new and independent evidence for the C—H bond dissociation energies in propylene. The existing thermochemical data lead to the following relationship between the C—H bond dissociation energy in propylene and the C—C bond dissociation energy in butene-1:

$$D(\text{CH}_2:\text{CH}.\text{CH}_2-\text{H}) - D(\text{CH}_2:\text{CH}.\text{CH}_2-\text{CH}_3) = 15.3 \text{ kcal./mole}.$$

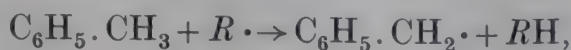
With a view to exploiting this relationship we have investigated the thermal decomposition of butene-1.

There seems to be no doubt that the weakest bond in butene-1 is



and in consequence the first step in the decomposition of this compound should yield methyl and allyl radicals. Although several studies of the pyrolysis of butene-1 are reported in the literature (Calingaert 1923; Hurd & Goldsby 1934), it is difficult

to formulate an unambiguous reaction mechanism because of complex side reactions between the radicals produced and the parent molecules. In order to eliminate these complicating factors we applied in the investigation of the thermal decomposition of butene-1 the technique essentially similar to that used for the estimation of the C—C bond dissociation energy in ethyl benzene (Szwarc 1949*c*). The main feature of this technique is the use of toluene as a carrier gas. The radicals formed in the decomposition of the compound under investigation are readily removed by the fast reaction with toluene



the stable benzyl radicals eventually dimerizing to dibenzyl. Thus in the particular case of butene-1 the initially formed methyl radicals are removed and the rate of the primary decomposition process is measured by the rate of formation of methane.

EXPERIMENTAL

Materials used

Butene-1 was generously supplied by Petrocarbon Ltd., Manchester, in the original cylinder as obtained from The Matheson Co., New Jersey, U.S.A. It was further purified by vacuum distillation from -80°C into a trap cooled by liquid air and connected to a mercury diffusion pump.

Toluene was prepared by repeated pyrolyses of 'Sulphur free' toluene at 810 to 820°C followed by a careful distillation of the pyrolyzed material.

Apparatus and technique

The diagram of the apparatus is given in figure 1. The storage flask *B* connected to manometer M_1 contained gaseous butene-1. This could be introduced through the needle valve *NV* and capillary K_1 into a large excess of toluene vapour, the total amount of butene-1 consumed during a run being given by the difference of pressure on M_1 .* The rate of flow of butene-1 could be varied by adjusting the needle valve as required. The toluene was introduced from flask *T* kept at constant temperature by means of a water bath *W*, the amount consumed being given by the difference in weight of flask *T*. The partial pressure of toluene could be varied by changing the temperature of bath *W*. The mixture of both compounds, maintained under a pressure of a few mm. Hg measured by means of a cathetometer on manometer M_2 , was made to flow through a cylindrical silica reaction vessel *R* heated by an electric furnace. The gases leaving the reaction vessel passed through a Pyrex tube and capillary K_2 (both heated electrically to about 80°C) into a trap *U* cooled by ice. The non-volatile product, namely, dibenzyl, was deposited in this trap. The undecomposed toluene and butene, and the products of decomposition other than methane and hydrogen were condensed in a large trap *H* at -78°C and a small trap *L* at -180°C . The non-condensed hydrogen and methane were continuously removed by a system of two mercury diffusion pumps P_1 and P_2 , and pumped through another trap

* For purposes of calculation it was assumed that butene-1 obeys the ideal gas law.

E cooled to -180°C into the storage section *S*. The amount of these gases could be estimated by measuring their pressure on a McLeod gauge. Subsequently part of these gases was admitted into an analytical system consisting of a copper oxide combustion cell and a pressure measurement device. This enabled us to estimate the ratio hydrogen to methane in a mixture of the two. The time of contact could be varied by sealing suitable capillaries K_2 into the tube leading from the reaction vessel to trap *U*.

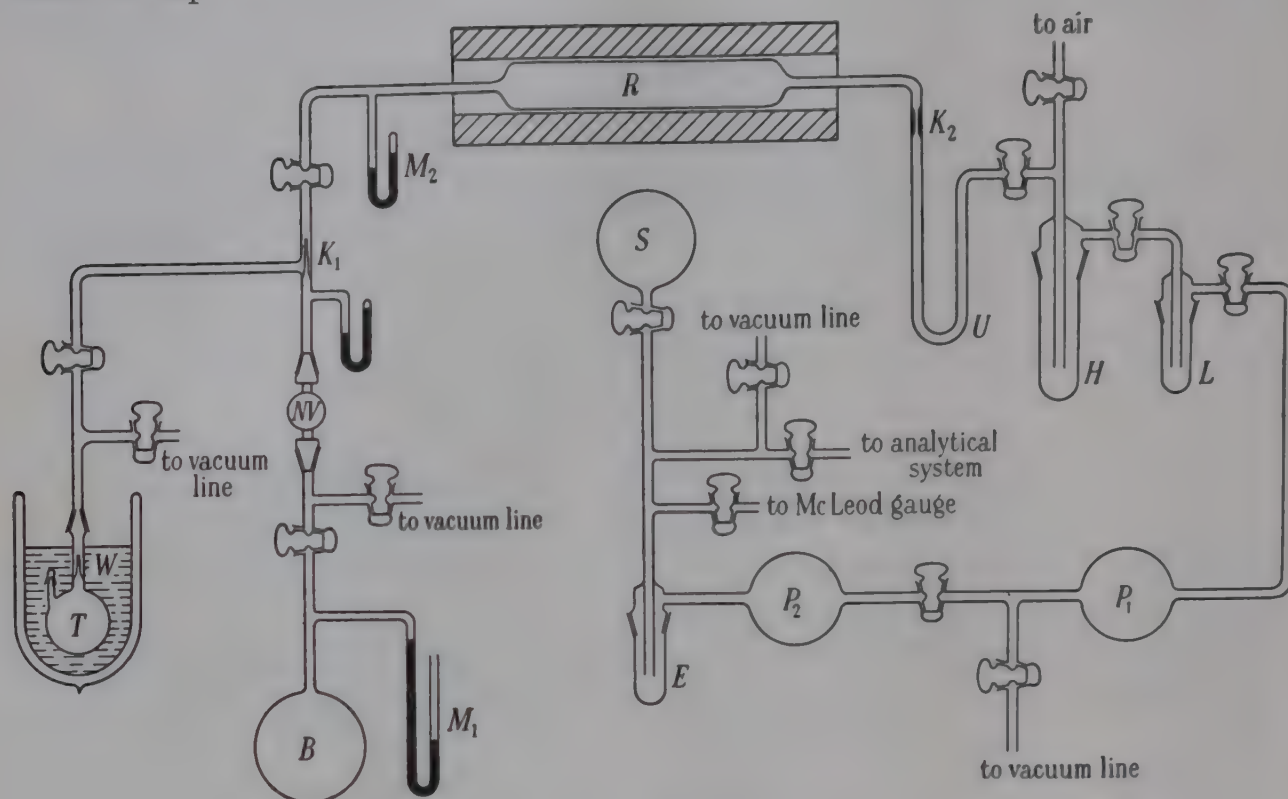


FIGURE 1. Apparatus.

Results

The pyrolysis of butene-1 was investigated over the temperature range 926 to 1046°K . Before each experiment a blank run for the pyrolysis of toluene only was carried out, and the appropriate corrections for its decomposition were introduced. This correction amounted to about 2% at the lower temperatures and to about 9% at the highest temperatures of pyrolysis.

The analysis of the non-condensable gases, summarized in table 1, showed that in addition to methane varying amounts of hydrogen were formed. The hydrogen content increased up to 25% with increase in temperature. The presence of hydrogen suggested a side reaction involving the direct dehydrogenation of butene-1 to butadiene; and indeed, this compound was detected by a semi-quantitative ultra-violet analysis of the gases condensed in the traps *H* and *L*. This analysis indicated that the percentage of butadiene increased with temperature.

The solid product accumulated in the *U* trap was proved to be dibenzyl, but the quantities obtained were definitely lower than the expected ones on the basis of a 1:1 molar ratio of methane to dibenzyl. The experimental results are given in table 2.

The decomposition of butene-1 measured by the rate of formation of the non-condensed gases was proved to be a homogeneous gas reaction of the first order. The variation of the surface/volume ratio by a factor of 4.5* had no effect on the first order rate constant, as can be seen from table 3. The change of the partial pressure of butene-1 by a factor of 6, as shown in table 4, and of the time of contact by a factor of 3, as shown in table 5, had no influence on the first order rate constant.

TABLE 1

run	T ($^{\circ}\text{K}$)	% H_2	% CH_4
38	935	3	97
59 <i>P</i>	964	8	92
52	993	12	88
51	996	16	84
58 <i>P</i>	998	13	87
57 <i>P</i>	1002	13.5	86.5
31	1009	15	85
54	1011	9	91
53	1014	8	92
29	1017	19	81
32	1027	21	79
33	1034	22	78
60 <i>P</i>	1035	24	76
34	1046	26	74
61 <i>P</i>	1051	26	74

The runs denoted by *P* refer to packed reaction vessel. Surface to volume ratio was increased by a factor of 4.5.

TABLE 2

runs	T ($^{\circ}\text{K}$)	$\frac{\text{millimols dibenzyl}}{\text{millimols } (\text{CH}_4 + \text{H}_2)}$
7, 8	about 996	1 : 2.7
9, 10		1 : 2.5
11, 12		1 : 2.5
13, 14, 15		1 : 3.5
16, 17		1 : 4.4
18, 19		1 : 2.5
20, 21		1 : 3.9
22, 23		1 : 2.4
24, 25		1 : 2.7
28, 29	about 1018	1 : 1.5
32, 33	about 1030	1 : 2.0
26, 27	about 1043	1 : 3.0

The variation of the partial pressure of toluene from 4 to 9 mm. Hg produced no change in the rate constant of the decomposition carried out in the lower temperature range (i.e. 954°K). At higher temperatures, however, as can be seen from table 6, the variation of the 'standard' total pressure of about 9 mm. Hg by a factor of 2 in both directions produced a change of about 30 % in the observed rate constant.

* The surface/volume ratio was changed by packing the reaction vessel with silica wool.

TABLE 3

run	T ($^{\circ}$ K)	$10^2 k$ (sec. $^{-1}$)	$10^2 k_{964^{\circ}}$ (sec. $^{-1}$)
T about 964° K			
5	962	15	16
59 <i>P</i>	964	17	17
40	966	17	16
T about 996° K			
			$10^2 k_{996^{\circ}}$ (sec. $^{-1}$)
14	997	42	41
42	997	47	45
17	998	42	39
58 <i>P</i>	998	47	44
43	999	46	42
41	1001	52	44
57 <i>P</i>	1002	53.5	44
T about 1035° K			
			$10^2 k_{1035^{\circ}}$ (sec. $^{-1}$)
33 <i>P</i>	1034	118	122
60	1035	121	121

Runs denoted by *P* correspond to the packed reaction vessel (surface increased by about 4.5).

TABLE 4

run	T ($^{\circ}$ K)	pressure of butene (mm. Hg)	$10^2 k$ (sec. $^{-1}$)	$10^2 k_{996^{\circ}}$ (sec. $^{-1}$)
18	993	0.104	39	43
12	993	0.171	41	45
11	996	0.214	41	41
8	997	0.220	42.5	41
14	997	0.327	42.5	41
13	994	0.338	39.5	42
15	996	0.402	47	47
17	998	0.580	42	39
16	996	0.593	39	39

TABLE 5

run	T ($^{\circ}$ K)	time of contact (sec.)	$10^2 k$ (sec. $^{-1}$)	$10^2 k_{993^{\circ}}$ (sec. $^{-1}$)
T about 993° K				
52	993	0.286	39	39
18	993	0.410	39	39
12	993	0.414	41	41
T about 1014° K				
				$10^2 k_{1014^{\circ}}$ (sec. $^{-1}$)
53	1014	0.156	81.5	81.5
54	1011	0.164	67.5	74
29	1017	0.362	76	69
28	1019	0.368	73	63
31	1009	0.426	63	73.5
30	1017	0.426	65	59
45	1019	0.437	74	63

The variation of the rate constant over the whole temperature range 926 to 1046°K is given in table 7. The plot of $\log k$ against $1/T$ is given in figure 2, the full line (a) corresponding to an activation energy of 63 kcal./mole, and a frequency factor of $2.5 \times 10^{13} \text{ sec.}^{-1}$.

TABLE 6

run	T ($^\circ\text{K}$)	pressure of toluene mm. Hg	$10^2 k$ (sec.^{-1})	$10^2 k_{964^\circ}$ (sec.^{-1})
T about 954°K				
50	954	4.80	10.6	10.6
39	970	9.09	18.4	10.7
40	966	9.14	17.1	11.3
T about 996°K				
				$10^2 k_{996^\circ}$ (sec.^{-1})
47	996	4.05	29.5	29.5
46	992	4.06	26	29
44	1005	4.19	41	31
20	992	4.96	38	43
21	994	4.96	40.5	43
24	995	7.55	37	38
8	997	8.66	42.5	41
18	993	8.88	39	43
11	996	8.91	41	41
12	993	8.99	41	45
7	996	8.99	42	42
43	999	9.41	46	42
13	994	9.42	39.5	42
42	997	9.44	47	45
14	997	9.46	42.5	41
16	996	9.56	39	39
17	998	9.60	42	39
51	996	9.94	46	46
25	996	13.7	51.5	51.5
22	995	15.3	54	56
23	996	15.5	55	55
T about 1044°K				
				$10^2 k_{1044^\circ}$ (sec.^{-1})
49	1043	4.68	136	140
48	1044	4.75	138	138
33	1034	9.20	118	158
26	1044	9.23	170	170
27	1042	9.34	178	190
34	1046	9.36	183	172

DISCUSSION

The above experimental results seem to indicate that the first step of the thermal decomposition of butene-1 is represented by reaction (5)



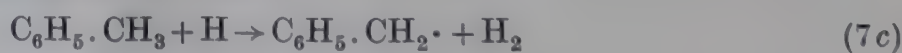
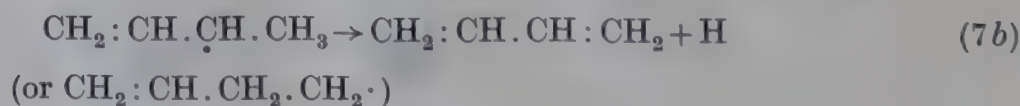
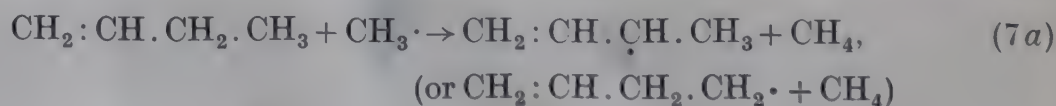
followed by the fast reaction (6)



TABLE 7

run	<i>T</i> (° K)	total pressure (mm. Hg)	pressure of butene (mm. Hg)	time of contact (sec.)	percent of decomp.	10 ² <i>k</i> (sec. ⁻¹)
35	926	9.26	0.262	0.466	1.8	3.9
36	927	9.35	0.283	0.468	1.88	4.0
37	929	9.18	0.218	0.457	2.16	4.8
38	935	9.38	0.306	0.460	2.62	5.8
3	941	9.42	0.390	0.528	2.66	5.1
4	962	8.85	0.212	0.495	6.3	13.3
5	962	8.13	0.135	0.460	6.7	15.0
59 <i>P</i>	964	9.88	0.224	0.494	8.2	17.3
2	965	9.61	0.695	0.486	5.5	11.2
40	966	9.14	0.276	0.439	7.25	17.1
39	970	9.09	0.226	0.438	8.0	18.4
18	993	8.88	0.104	0.410	14.7	39
12	993	8.99	0.171	0.414	15.6	41
52	993	9.39	0.177	0.286	10.5	39
13	994	9.42	0.338	0.431	15.9	39
11	996	8.91	0.214	0.410	15.4	41
7	996	8.99	0.210	0.418	16.0	42
15	996	9.39	0.402	0.431	18.2	47
16	996	9.56	0.593	0.416	14.9	39
51	996	9.94	0.272	0.424	17.8	46.0
8	997	8.66	0.220	0.401	15.6	42
42	997	9.44	0.229	0.419	17.8	47
14	997	9.46	0.327	0.439	17.0	42
17	998	9.60	0.580	0.416	16.0	42
58 <i>P</i>	998	9.69	0.252	0.462	20.5	47
43	999	9.41	0.286	0.438	18.3	46
41	1001	9.50	0.184	0.425	19.8	52
57 <i>P</i>	1002	9.72	0.230	0.462	22.0	53
31	1009	9.39	0.297	0.426	23.5	63
54	1011	9.13	0.094	0.164	10.5	67
53	1014	9.09	0.101	0.156	11.9	81
29	1017	8.18	0.255	0.362	24.1	76
30	1017	9.44	0.302	0.426	24.2	65
28	1019	8.28	0.224	0.368	23.6	73
45	1019	9.24	0.266	0.437	27.5	74
32	1027	9.34	0.273	0.420	34	99
33	1034	9.20	0.302	0.408	38	118
60 <i>P</i>	1035	9.87	0.286	0.454	42	121
55	1040	9.02	0.258	0.413	47	154
27	1042	9.34	0.254	0.402	51	178
26	1044	9.23	0.262	0.397	49	170
34	1046	9.36	0.266	0.405	52	183

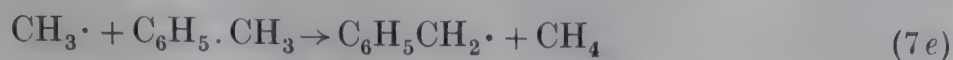
producing methane. The observed hydrogen cannot be entirely accounted for by reactions



and



followed by



The molar ratio of butene-1 to toluene, which was 1 : 30 in most of the experiments, makes the probability of the reaction (7a) to be about 1/30 of the probability of reaction (6). The maximum amount of hydrogen atoms produced by reactions (7a) and (7b) can be, therefore, about 3 % of the amount of methane observed. Since the

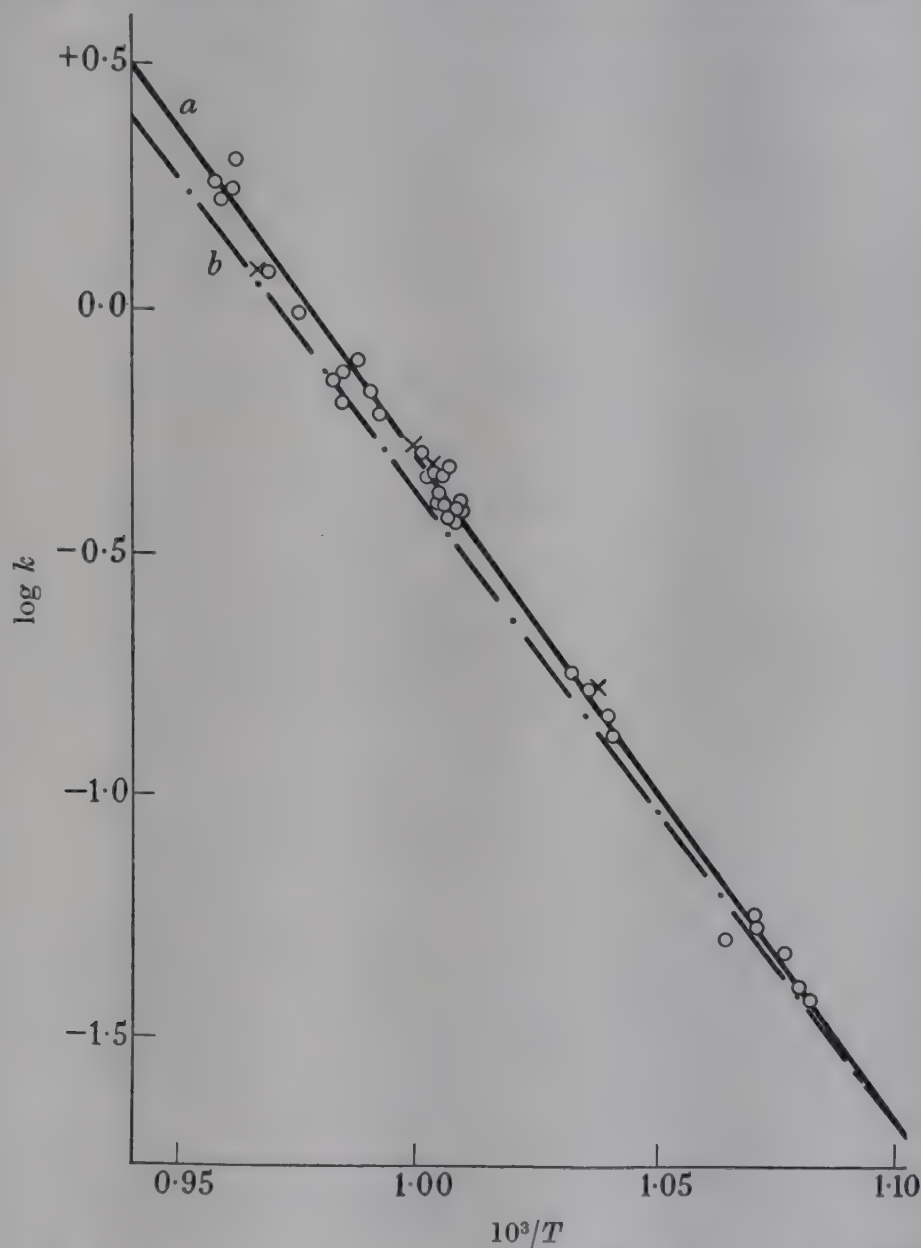


FIGURE 2. Curve *a*, $E = 63$ kcal./mole, $\nu = 2.5 \times 10^{13}$ sec. $^{-1}$; *b*, $E = 60$ kcal./mole, $\nu = 5.0 \times 10^{12}$ sec. $^{-1}$. $\times$, packed reaction vessel; $\circ$, non-packed reaction vessel.

reaction between hydrogen atoms and toluene (Szwarc 1948) produces hydrogen and methane in the ratio 3 : 2 the maximum quantity of hydrogen molecules accounted for by reactions (7a), (7b), (7c), (7d) can amount to 2 per cent only of the observed quantity. It seems plausible, therefore, to attribute the main quantity of hydrogen and butadiene produced to the side reaction



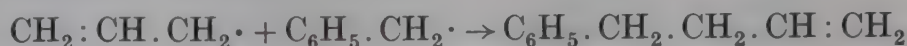
It could be suggested that hydrogen was produced by the decomposition of allyl radicals into allene and hydrogen atoms



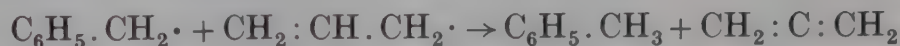
followed by reactions (7c), (7d). If this were the case, then we should obtain one mole of dibenzyl for each mole of gas ($\text{CH}_4 + \text{H}_2$) formed. Inspection of table 2 shows, however, that the ratio of dibenzyl/($\text{CH}_4 + \text{H}_2$) is much smaller than 1:1 being occasionally as low as 1:4. This observation can be explained by the stability of allyl radicals under our experimental conditions. It follows that allyl radicals do not react readily with toluene according to the equation



and that a considerable fraction of these is removed instead either by a reaction between themselves or by reaction with benzyl radicals. This last reaction may lead either to the association product



or to disproportionation products



Both reactions would account for the decrease in the observed amounts of dibenzyl. These results are supported by an analogous observation made in the investigation of the pyrolysis of allyl bromide (Szwarc & Ghosh 1949) the amounts of dibenzyl formed in this case being also considerably lower than the expected values.

Activation energies

In order to determine the activation energy of process (5) we have adopted three methods:

(i) The plot of $\log k$ against $1/T$ is shown in figure 2 by the full line (a); k denotes here the unimolecular rate constant for the overall decomposition calculated on the assumption that the rate of decomposition is measured by the rate of formation of non-condensed gases. This line gives the activation energy as 63 kcal./mole and the frequency factor $2.5 \times 10^{13} \text{ sec.}^{-1}$.

(ii) By making the appropriate correction for the presence of hydrogen, the dotted line (b) is obtained. The activation energy is then 60 kcal./mole only and the frequency factor $5 \times 10^{12} \text{ sec.}^{-1}$.

(iii) Assuming the value of $10^{13} \text{ sec.}^{-1}$ for the frequency factor of reaction (5) and using the Arrhenius equation $k = 10^{13} e^{-E/RT} \text{ sec.}^{-1}$ where k is the rate constant based on the formation of methane only, we arrived at the average value of 61.5 kcal./mole for the activation energy (see table 8).

We favour this latter value for the activation energy on the following grounds:

(a) This value was based on results obtained in the lower temperature region where the extent of the disturbing side reactions producing hydrogen was negligible.

(b) In this series of experiments the rate constant was unaffected by the variation of toluene pressure. This indicates that the decomposition was a truly unimolecular

process. On the other hand the results obtained at higher temperatures exhibited some dependence on the pressure of toluene.* A similar dependence on toluene pressure was also observed both in the case of the decomposition of hydrazine (Szwarc 1949*d*) and in the case the decomposition of methyl bromide carried out recently in this laboratory. In our opinion this effect is due to a decrease in the stationary concentration of the energized molecules. In the high temperature region the rate of decomposition becomes too rapid and the rate of energy transfer under our experimental conditions is not sufficiently high to maintain the normal Maxwell distribution of the energized molecules.† In consequence the activation energy obtained from the slope of line (b) is lower than the 'true' activation energy.

TABLE 8

run	$T(^{\circ}\text{K})$	$10^2 k_{\text{CH}_3}$ (sec^{-1})	E kcal./mole
38	935	5.6	61.4
50	954	10.1	61.4
59 <i>P</i>	964	15.9	61.1
52	993	33.9	61.6
47	996	26.8	62.3
51	996	38.0	61.5
58 <i>P</i>	998	42.4	61.4

Bond dissociation energies

Making the usual assumption that the recombination of radicals formed does not involve any activation energy, we conclude that the dissociation energy of the 3.4 C—C bond in butene-1 is 61.5 kcal./mole. This value leads to

$$D(\text{CH}_2:\text{CH}.\text{CH}_2-\text{H}) = 76.5 \text{ kcal./mole,}$$

to about 40 kcal./mole for the dissociation energy of the central C—C bond in diallyl, and to +30 kcal./mole for the heat of formation of the allyl radical.

Furthermore, the above value for $D(\text{C—H})$ in propylene leads to the value of 68.5 kcal./mole for the C—H bond dissociation energy in the allyl radical. It is interesting to note that this C—H bond dissociation energy is higher by 7 kcal./mole than the C—C bond dissociation energy in butene-1. This is an additional argument showing that under our experimental conditions the allyl radicals produced do not decompose further.

Assuming that the decrease of the C—H bond dissociation energy in propylene as compared with the C—H bond dissociation energy in methane is attributed to the resonance energy of the allyl radical we obtain for the latter the value of 25.5 kcal./mole.

On similar lines it was shown (Szwarc 1947, 1948) that the benzyl radical is stabilized by 24.5 kcal./mole. It is necessary to mention, however, that although both radicals are almost equally stabilized by their respective resonance energies, the

* Technical reasons prevented us from investigating the decomposition at pressure of toluene higher than 16 mm. Hg.

† The problem of energy transfer is now under investigation and further results will be reported in due course.

thermal stability of the benzyl radical is much higher than that of the allyl radical. Whereas the benzyl radical does not decompose into a still more stable entity, the allyl radical may decompose into allene and hydrogen atom. From the previous discussion it follows that the endothermicity of this decomposition is only 68.5 kcal./mole. In consequence propylene would be less suitable than toluene as a 'radical remover' in the pyrolytic decompositions at higher temperatures.

We are indebted to Petrocarbon Ltd. for the supply of Butene-1 and for the ultra-violet analyses, and to Monsanto Chemicals Ltd. for the research grant to one of us (A.H.S.).

The authors wish to express their thanks to Professor M. G. Evans, F.R.S., for his constant interest and encouragement during the course of this work.

REFERENCES

- Butler, E. T. & Polanyi, M. 1943 *Trans. Faraday Soc.* **39**, 19.
Calingaert, G. 1923 *J. Amer. Chem. Soc.* **45**, 130.
Coulson, C. A. 1938 *Proc. Roy. Soc. A*, **164**, 383.
Hurd, C. D. & Goldsby, A. R. 1934 *J. Amer. Chem. Soc.* **56**, 1812.
Orr, W. J. C. 1946 (quoted by Bolland, J. L. & Gee, G.) *Trans. Faraday Soc.* **42**, 244.
Rice, F. O. & Polly, O. L. 1938 *J. Chem. Phys.* **6**, 273.
Smith, J. R. E. & Hinshelwood, C. N. 1942 *Proc. Roy. Soc. A*, **180**, 237.
Szwarc, M. 1947 *Nature* **160**, 403.
Szwarc, M. 1948 *J. Chem. Phys.* **16**, 128.
Szwarc, M. 1949a *J. Chem. Phys.* **17**, 284.
Szwarc, M. 1949b *J. Chem. Phys.* **17**, 292.
Szwarc, M. 1949c *J. Chem. Phys.* **17**, 431.
Szwarc, M. 1949d *Proc. Roy. Soc. A*, **198**, 267.
Szwarc, M. & Ghosh, B. N. 1949 *J. Chem. Phys.* **17**, 744.
Taylor, H. S. & Smith, J. O. 1940 *J. Chem. Phys.* **8**, 543.

Diffraction by a plane screen

BY E. T. COPSON

(Communicated by Sir Edmund Whittaker, F.R.S.—

Received 20 March 1950)

The integral-equation method of solving the problem of the diffraction of electromagnetic waves by a perfectly conducting plane screen has been criticized by C. J. Bouwkamp, who claims that it is valid only when certain boundary conditions are satisfied on the edge of the screen. This criticism is answered. It is also shown that, since the equations to be solved are differential-integral equations, an arbitrary function arises in the solution and that this arbitrary function may be chosen so that, although there are singularities at the edge of the screen, there is no radiation of energy from the edge. As an illustration, the three-dimensional problem of diffraction by a half-plane is solved.

1. THE INTEGRAL EQUATIONS OF DIFFRACTION BY A PLANE SCREEN

In a previous paper (Copson 1946*b*), the problem of the diffraction of electromagnetic waves was proved to depend on the solution of a system of integral equations, and the method was applied to Sommerfeld's two-dimensional problem of diffraction by a half-plane. It was also shown that the method led to the approximate solutions of certain three-dimensional problems discussed by Rayleigh and Bethe. Integral-equation methods have also been applied in America by Schwinger and his collaborators to problems of diffraction by more complicated screens, and by Fox to the problem of the diffraction of sound pulses (Fox 1948, 1949).

Despite this, the method has been criticized in a review (Bouwkamp 1947). The point at issue can be illustrated by reference to Theorem 5 of my paper (Copson 1946*b*, p. 109). If an electromagnetic field* $\mathbf{E}^i, \mathbf{H}^i$ is incident in the half-space $z > 0$ on a perfectly conducting screen in the plane $z = 0$, the scattered field $\mathbf{E}^s, \mathbf{H}^s$ is expressed in the form

$$\begin{aligned} E_x^s &= ikv' - \frac{\partial w'}{\partial x}, & E_y^s &= -iku' - \frac{\partial w'}{\partial y}, & E_z^s &= -\frac{\partial w'}{\partial z}, \\ H_x^s &= -\frac{\partial u'}{\partial z}, & H_y^s &= -\frac{\partial v'}{\partial z}, & H_z^s &= \frac{\partial u'}{\partial x} + \frac{\partial v'}{\partial y}, \end{aligned}$$

where u', v', w' are three solutions of the partial differential equation $(\nabla^2 + k^2)\Omega = 0$. In order that the scattered field may satisfy Maxwell's equations, it is necessary that

$$ikw' = \frac{\partial v'}{\partial x} - \frac{\partial u'}{\partial y}. \quad (1.1)$$

The three wave functions are given explicitly by

$$(u', v', w') = \frac{1}{2\pi} \iint_{S_1} (h_x, h_y, e_z) \phi dx' dy', \quad (1.2)$$

* A time factor $e^{i\omega t}$ ($\omega > 0$) is to be understood throughout this paper, which is concerned only with monochromatic waves.

where S_2 is the metal of the screen and $\phi = e^{-ikR}/R$ with

$$R^2 = (x-x')^2 + (y-y')^2 + z^2.$$

In (1.2), h_x, h_y and e_z are the unknown values of H_x^s, H_y^s, E_z^s on the positive face of S_2 , and, in virtue of Maxwell's equations, are connected by

$$ike_z = \frac{\partial h_y}{\partial x'} - \frac{\partial h_x}{\partial y'}. \quad (1.3)$$

Bouwkamp's objection is that the functions defined by (1.2) and (1.3) do not, in general, satisfy (1.1).

His proof of this is as follows. Differentiation under the sign of integration followed by an application of Green's theorem gives

$$\begin{aligned} \frac{\partial v'}{\partial x} - \frac{\partial u'}{\partial y} &= -\frac{1}{2\pi} \iint_{S_1} \left\{ h_y \frac{\partial \phi}{\partial x'} - h_x \frac{\partial \phi}{\partial y'} \right\} dx' dy' \\ &= \frac{1}{2\pi} \iint_{S_1} \left(\frac{\partial h_y}{\partial x'} - \frac{\partial h_x}{\partial y'} \right) \phi dx' dy' - \frac{1}{2\pi} \int_{\Gamma} \phi (h_x dx' + h_y dy') \end{aligned}$$

where Γ is the rim of S_2 . Hence (1.1) is a consequence of (1.2) and (1.3) only if the line integral along Γ vanishes. A similar criticism applies to Theorem 4 (Copson 1946*b*, p. 106).

It happens to be the case that the conditions were satisfied in all the applications of these theorems, but it is not necessarily always so. The satisfactory results given by the integral-equation method have led to the conjecture that 'in so far as the fields at a distance from the screen are concerned, it does not appear that the boundary condition on the rim is of any great importance' (Miles 1949, p. 763). This is by no means the case, as a singularity of too high an order on the rim completely alters the character of the solution (see Rayleigh 1897; Rayleigh 1913; Bouwkamp 1946; Baker & Copson 1950, p. 185).

The solution of this difficulty is actually quite simple. Bouwkamp's argument depends on differentiation under the sign of integration and the use of Green's theorem, and is valid only when the line integral vanishes. For h_x is the value of H_x^s in the plane $z = 0$ and vanishes in the aperture. If h_x is continuous, it vanishes on the rim of the screen, whereas, if it is discontinuous, it becomes infinite on the rim, as Bouwkamp has himself pointed out (Bouwkamp 1946). And similarly for h_y . Thus either h_x and h_y vanish on the rim Γ , and (1.1) is a consequence of (1.2) and (1.3); or else at least one of the function h_x and h_y is infinite on Γ , and Bouwkamp's argument fails. Similar remarks apply to Theorem 4.

An examination of the proof of Theorem 5 shows that it is not necessary to introduce the function w' at all, and that if the theorem is expressed in a form involving only u' and v' , Bouwkamp's criticism is not relevant. And similarly for Theorem 4. The modified forms of these theorems are as follows.

THEOREM A. *Let an electromagnetic field $\mathbf{E}^i, \mathbf{H}^i$ be incident in the half-space $z > 0$ on a perfectly conducting screen in the plane $z = 0$, the holes in the screen being*

denoted by S_1 , the metal of the screen by S_2 . Then the total field in the half-space $z < 0$ is

$$\left. \begin{aligned} E_x &= \frac{\partial u}{\partial z}, & H_x &= ikv - \frac{1}{ik} \frac{\partial^2 v}{\partial x^2} + \frac{1}{ik} \frac{\partial^2 u}{\partial x \partial y}, \\ E_y &= \frac{\partial v}{\partial z}, & H_y &= -iku + \frac{1}{ik} \frac{\partial^2 u}{\partial y^2} - \frac{1}{ik} \frac{\partial^2 v}{\partial x \partial y}, \\ E_z &= -\frac{\partial u}{\partial x} - \frac{\partial v}{\partial y}, & H_z &= \frac{1}{ik} \frac{\partial^2 u}{\partial y \partial z} - \frac{1}{ik} \frac{\partial^2 v}{\partial x \partial z}, \end{aligned} \right\} \quad (1.4)$$

where

$$(u, v) = \frac{1}{2\pi} \iint_{S_1} (e_x, e_y) \phi dx' dy'. \quad (1.5)$$

The functions e_x, e_y satisfy the equations

$$\left. \begin{aligned} \frac{\partial^2 u_0}{\partial x^2} + \frac{\partial^2 u_0}{\partial y^2} + k^2 u_0 &= -\left(\frac{\partial E_x^i}{\partial z}\right)_{z=0}, \\ \frac{\partial^2 v_0}{\partial x^2} + \frac{\partial^2 v_0}{\partial y^2} + k^2 v_0 &= -\left(\frac{\partial E_y^i}{\partial z}\right)_{z=0}, \\ \frac{\partial u_0}{\partial x} + \frac{\partial v_0}{\partial y} &= -(E_z^i)_{z=0}, \end{aligned} \right\} \quad (1.6)$$

where u_0, v_0 are the values of u, v at a point $(x, y, 0)$ of S_1 .

It should be noted that equations (1.6) are consistent. For since $\text{div } \mathbf{E}^i = 0$ and $(\nabla^2 + k^2) E_z^i = 0$, it follows from the first two that

$$\left(\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + k^2\right) \left(\frac{\partial u_0}{\partial x} + \frac{\partial v_0}{\partial y}\right) = -\left(\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + k^2\right) (E_z^i)_{z=0}.$$

THEOREM B. Under the conditions of theorem A, the total field everywhere is

$$\left. \begin{aligned} E_x &= E_x^i + ikv' - \frac{1}{ik} \frac{\partial^2 v'}{\partial x^2} + \frac{1}{ik} \frac{\partial^2 u'}{\partial x \partial y}, & H_x &= H_x^i - \frac{\partial u'}{\partial z}, \\ E_y &= E_y^i - iku' + \frac{1}{ik} \frac{\partial^2 u'}{\partial y^2} - \frac{1}{ik} \frac{\partial^2 v'}{\partial x \partial y}, & H_y &= H_y^i - \frac{\partial v'}{\partial z}, \\ E_z &= E_z^i + \frac{1}{ik} \frac{\partial^2 u'}{\partial y \partial z} - \frac{1}{ik} \frac{\partial^2 v'}{\partial x \partial y}, & H_z &= H_z^i + \frac{\partial u'}{\partial x} + \frac{\partial v'}{\partial y}, \end{aligned} \right\} \quad (1.7)$$

where

$$(u', v') = \frac{1}{2\pi} \iint_{S_1} (h_x, h_y) \phi dx' dy'. \quad (1.8)$$

The functions h_x, h_y satisfy the equations

$$\left. \begin{aligned} \frac{\partial^2 u'_0}{\partial x^2} + \frac{\partial^2 u'_0}{\partial y^2} + k^2 u'_0 &= -\left(\frac{\partial H_x^i}{\partial z}\right)_{z=0}, \\ \frac{\partial^2 v'_0}{\partial x^2} + \frac{\partial^2 v'_0}{\partial y^2} + k^2 v'_0 &= -\left(\frac{\partial H_y^i}{\partial z}\right)_{z=0}, \\ \frac{\partial u'_0}{\partial x} + \frac{\partial v'_0}{\partial y} &= -(H_z^i)_{z=0}, \end{aligned} \right\} \quad (1.9)$$

where u'_0, v'_0 are the values of u', v' at a point $(x, y, 0)$ of S_2 .

2. SINGULARITIES ON THE RIM OF THE SCREEN

It has been pointed out (Bouwkamp 1946) that, in the solution of the problem of diffraction of electromagnetic waves, there may arise wave functions which are infinite on the rim of the screen, and that, if this is so, the order of the singularity must be such that no energy is radiated from the apparent sources on the rim. The condition is satisfied if no component of $\mathbf{E}$ or $\mathbf{H}$ becomes infinite to a higher order than ρ^{-1} , where ρ is the distance from the rim, as is indeed the case in Sommerfeld's solution of the two-dimensional cases of diffraction of plane waves by a half-plane. It seems likely that the convergence of the integrals defining u and v in theorem A, or u' and v' in theorem B would suffice, but a rigorous proof would be difficult.

Bouwkamp also points out that differentiation of a wave function which is infinite on the rim increases the order of the singularity; and since the formulae of theorems A and B involve second derivatives, they would therefore be suspect. This is, however, not necessarily the case. Let us consider theorem A. The function u_0 satisfies in the domain S_1 the partial differential equation

$$\frac{\partial^2 u_0}{\partial x^2} + \frac{\partial^2 u_0}{\partial y^2} + k^2 u_0 = - \left(\frac{\partial E_x^i}{\partial z} \right)_{z=0},$$

which is of elliptic type. When this equation is solved by the Green function method, the boundary values on Γ either of u_0 or of its normal derivative can be assigned arbitrarily; and similarly for v_0 . Thus the solution of the two second-order partial differential equations in (1.6) involves two arbitrary functions, between which a relation is set up by the remaining equation

$$\frac{\partial u_0}{\partial x} + \frac{\partial v_0}{\partial y} = - (E_z^i)_{z=0}.$$

Hence the functions u_0 and v_0 depend on one arbitrary function which can be used to prevent the occurrence of singularities of too high an order. And similarly for theorem B. Again, a rigorous proof of this would be difficult, but, in the absence of an existence theorem, it is necessary to rely here on physical intuition.

The way in which the arbitrary function can be chosen to prevent the occurrence of singularities of too high an order is well illustrated by the three-dimensional case of diffraction by a half-plane which has recently been solved by P. C. Clemmow.*

3. DIFFRACTION BY A HALF-PLANE

In the three-dimensional problem, the field

$$\mathbf{E}^i = (-m, l, 0) e^{ik(lx+my+nz)}, \quad \mathbf{H}^i = (ln, mn, -l^2 - m^2) e^{ik(lx+my+nz)}$$

is incident on the perfectly conducting screen $z = 0$, $x > 0$. Here (l, m, n) are the direction cosines of the incident waves so that $l^2 + m^2 + n^2 = 1$, and we assume, as in

* I must thank Mr Clemmow for drawing my attention to this problem and for an interesting critical discussion. His solution (Clemmow 1950) consists in observing that, if the edge of the screen is the axis of y , the variable y can occur only in the form e^{ikmy} , which enables him to reduce the problem to a two-dimensional one.

previous papers, that $k = p - iq$, where $q > 0$. The functions u'_0, v'_0 of theorem B satisfy equations (1.9) in the half-plane $x > 0$, and these equations reduce to

$$\frac{\partial^2 u'_0}{\partial x^2} + \frac{\partial^2 u'_0}{\partial y^2} + k^2 u'_0 = -ikln^2 e^{ik(lx+my)},$$

$$\frac{\partial^2 v'_0}{\partial x^2} + \frac{\partial^2 v'_0}{\partial y^2} + k^2 v'_0 = -ikmn^2 e^{ik(lx+my)},$$

$$\frac{\partial u_0}{\partial x} + \frac{\partial v_0}{\partial y} = (l^2 + m^2) e^{ik(lx+my)}.$$

Since y can occur only as a factor e^{ikmy} , it readily follows that

$$u'_0 = \frac{l}{ik} e^{ik(lx+my)} - mA e^{ik(\lambda x+my)} + mB e^{ik(-\lambda x+my)},$$

$$v'_0 = \frac{m}{ik} e^{ik(lx+my)} + \lambda A e^{ik(\lambda x+my)} + \lambda B e^{ik(-\lambda x+my)},$$

where $\lambda = +\sqrt{(l^2 + n^2)}$, and A and B are arbitrary constants. But since

$$|e^{\pm ik\lambda x}| = e^{\pm q\lambda x},$$

the constant A must be zero, for otherwise the term involving A would be infinite at infinity. The constant B , it will appear, is determined by the condition that no component of $\mathbf{E}$ or $\mathbf{H}$ is to be infinite on the edge Oy to too high an order. We have thus to solve the pair of integral equations

$$\frac{1}{2\pi} \iint_{S_1} h_x(x', y') \phi_0 dx' dy' = \frac{l}{ik} e^{ik(lx+my)} + mB e^{ik(-\lambda x+my)},$$

$$\frac{1}{2\pi} \iint_{S_1} h_y(x', y') \phi_0 dx' dy' = \frac{m}{ik} e^{ik(lx+my)} + \lambda B e^{ik(-\lambda x+my)},$$

when $x > 0$, the nucleus being $\phi_0 = \frac{e^{-ikr}}{r}$,

where $r^2 = (x - x')^2 + (y - y')^2$.

It follows that

$$\left. \begin{aligned} h_x(x', y') &= \frac{l}{ik} f(x', y') + mBg(x', y'), \\ h_y(x', y') &= \frac{m}{ik} f(x', y') + \lambda Bg(x', y'). \end{aligned} \right\} \quad (3.1)$$

where

$$\frac{1}{2\pi} \iint_{S_1} f(x', y') \phi_0 dx' dy' = e^{ik(lx+my)},$$

$$\frac{1}{2\pi} \iint_{S_1} g(x', y') \phi_0 dx' dy' = e^{ik(-\lambda x+my)},$$

when $x > 0$.

Evidently f and g can involve y' only as a factor $e^{ikmy'}$. If we substitute

$$f(x', y') = f(x') e^{ikmy'}, \quad g(x', y') = g(x') e^{ikmy'}, \quad (3.2)$$

and integrate with respect to y' , we obtain

$$\int_0^\infty f(x') H_0^{(2)}(\kappa |x-x'|) dx' = 2i e^{i\kappa x \cos \alpha},$$

$$\int_0^\infty g(x') H_0^{(2)}(\kappa |x-x'|) dx' = 2i e^{-i\kappa x}$$

where $k\lambda = \kappa$, $kl = \kappa \cos \alpha$, $kn = \kappa \sin \alpha$. The solution of the first integral equation is (cf. Copson, 1946*a*; Baker & Copson 1950, p. 168)

$$f(x) = i\kappa \sin \alpha e^{i\kappa x \cos \alpha} - \frac{1}{2\pi} \int_{-\infty}^\infty \frac{\sqrt{(\kappa - \kappa \cos \alpha)} \sqrt{(\kappa - w)}}{w + \kappa \cos \alpha} e^{-ixw} dw,$$

whereas the second, which corresponds to $\alpha = \pi$ in the first, has the simple solution

$$g(x) = \sqrt{\left(\frac{2i\kappa}{\pi x}\right)} e^{-i\kappa x}.$$

If we now substitute the values of h_x and h_y given by (3.1) and (3.2) in (1.8), we obtain

$$\left. \begin{aligned} u' &= \frac{l}{ik} \varpi e^{ikmy} + mB\varpi_0 e^{ikmy}, \\ v' &= \frac{m}{ik} \varpi e^{ikmy} + \lambda B\varpi_0 e^{ikmy}, \end{aligned} \right\} \quad (3.3)$$

where

$$\varpi = \frac{e^{-ikmy}}{2\pi} \iint_{S_1} f(x') e^{ikmy'} \phi dx' dy' = \frac{1}{2i} \int_0^\infty f(x') H_0^{(2)}[\kappa \sqrt{(x-x')^2 + z^2}] dx',$$

$$\varpi_0 = \frac{e^{-ikmy}}{2\pi} \iint_{S_1} g(x') e^{ikmy'} \phi dx' dy' = \frac{1}{2i} \int_0^\infty g(x') H_0^{(2)}[\kappa \sqrt{(x-x')^2 + z^2}] dx'.$$

It is convenient at this stage to introduce polar co-ordinates

$$x = \rho \cos \theta, \quad z = \rho \sin \theta,$$

so that ρ is the distance from the edge of the screen. The integral for ϖ can be transformed into Sommerfeld's expression

$$\begin{aligned} \varpi &= e^{i\kappa \rho \cos(\theta-\alpha)} - \frac{e^{\frac{1}{2}\pi i}}{\sqrt{\pi}} e^{i\kappa \rho \cos(\theta-\alpha)} \int_{-\infty}^{\sqrt{(2\kappa\rho) \cos \frac{1}{2}(\theta-\alpha)}} e^{-i\tau^2} d\tau \\ &\quad + \frac{e^{\frac{1}{2}\pi i}}{\sqrt{\pi}} e^{i\kappa \rho \cos(\theta+\alpha)} \int_{-\infty}^{\sqrt{(2\kappa\rho) \cos \frac{1}{2}(\theta+\alpha)}} e^{-i\tau^2} d\tau; \end{aligned}$$

and ϖ_0 simplifies to

$$\varpi_0 = 2 e^{-i\kappa \rho \cos \theta} - \frac{2 e^{\frac{1}{2}\pi i}}{\sqrt{\pi}} e^{-i\kappa \rho \cos \theta} \int_{-\infty}^{\sqrt{(2\kappa\rho) \sin \frac{1}{2}\theta}} e^{-i\tau^2} d\tau.$$

When $\kappa\rho$ is small, ϖ and ϖ_0 can be expanded as convergent series of powers of $\sqrt{(\kappa\rho)}$, the first two terms in the expansions being

$$\varpi = 1 - \frac{2 e^{\frac{1}{2}\pi i}}{\sqrt{\pi}} \sqrt{(2\kappa\rho)} \sin \frac{1}{2}\theta \sin \frac{1}{2}\alpha + \dots,$$

$$\varpi_0 = 1 - \frac{2 e^{\frac{1}{2}\pi i}}{\sqrt{\pi}} \sqrt{(2\kappa\rho)} \sin \frac{1}{2}\theta + \dots$$

Hence ϖ and ϖ_0 are equal to unity at the edge of the screen, but their first derivatives with respect to x and z become infinite like $\rho^{-\frac{1}{2}}$, their second derivatives like $\rho^{-\frac{3}{2}}$, and so on.

When the values of u' and v' given by (3.3) are substituted in (1.7), it turns out that second derivatives of ϖ and ϖ_0 occur only in the formulae for E_x and E_z , where there are terms $\partial^2 v' / \partial x^2$ and $\partial^2 v' / \partial x \partial z$. Unless B is properly chosen, these terms will make E_x and E_z infinite like $\rho^{-\frac{1}{2}}$ at the edge of the screen, and this is not allowed by Bouwkamp's criterion; but, with a proper choice of B , E_x and E_z are infinite only like $\rho^{-\frac{3}{2}}$. For it is readily proved that

$$\frac{\partial \varpi}{\partial x} = ikl\varpi + \frac{e^{\frac{1}{2}\pi i}}{\sqrt{\pi}} \sqrt{\left(\frac{2\kappa}{\rho}\right)} e^{-i\kappa\rho} \sin \frac{1}{2}\theta \sin \frac{1}{2}\alpha$$

$$\frac{\partial \varpi_0}{\partial x} = -i\kappa\varpi_0 + \frac{e^{\frac{1}{2}\pi i}}{\sqrt{\pi}} \sqrt{\left(\frac{2\kappa}{\rho}\right)} e^{-i\kappa\rho} \sin \frac{1}{2}\theta,$$

so that

$$\frac{\partial v'}{\partial x} = e^{ikmy} \left\{ lm\varpi - i\kappa\lambda B\varpi_0 + \lambda \left(B + \frac{m}{i\kappa} \sin \frac{1}{2}\alpha \right) \frac{e^{\frac{1}{2}\pi i}}{\sqrt{\pi}} \sqrt{\left(\frac{2\kappa}{\rho}\right)} e^{-i\kappa\rho} \sin \frac{1}{2}\theta \right\};$$

hence if

$$B = -\frac{m}{i\kappa} \sin \frac{1}{2}\alpha,$$

$\partial v' / \partial x$ is finite at the edge and E_x and E_z become infinite only like $\rho^{-\frac{3}{2}}$. The correct values of u' and v' are therefore

$$u' = \frac{e^{ikmy}}{ik} \left\{ l\varpi - \frac{m^2}{\lambda} \varpi_0 \sin \frac{1}{2}\alpha \right\},$$

$$v' = \frac{m e^{ikmy}}{ik} \{ \varpi - \varpi_0 \sin \frac{1}{2}\alpha \}.$$

Substitution of these values of u' and v' in (1.7) leads to the final result, that the total field everywhere is

$$E_x = - \left\{ m\Phi + m \left(\frac{2}{\pi\kappa\rho} \right)^{\frac{1}{2}} e^{-i\kappa\rho - \frac{1}{2}\pi i} \sin \frac{1}{2}\theta \sin \frac{1}{2}\alpha \right\} e^{ikmy},$$

$$E_y = l\Phi e^{ikmy},$$

$$E_z = m \left(\frac{2}{\pi\kappa\rho} \right)^{\frac{1}{2}} e^{-i\kappa\rho - \frac{1}{2}\pi i} \cos \frac{1}{2}\theta \sin \frac{1}{2}\alpha e^{ikmy},$$

$$H_x = \left\{ ln\Psi + (l\lambda - m^2) \left(\frac{2}{\pi\kappa\rho} \right)^{\frac{1}{2}} e^{-i\kappa\rho - \frac{1}{2}\pi i} \cos \frac{1}{2}\theta \sin \frac{1}{2}\alpha \right\} e^{ikmy},$$

$$H_y = mn\Psi e^{ikmy},$$

$$H_z = - \left\{ (l^2 + m^2) \Phi - (l\lambda - m^2) \left(\frac{2}{\pi\kappa\rho} \right)^{\frac{1}{2}} e^{-i\kappa\rho - \frac{1}{2}\pi i} \sin \frac{1}{2}\theta \sin \frac{1}{2}\alpha \right\} e^{ikmy},$$

where

$$\Phi = \frac{1}{\sqrt{\pi}} e^{i\kappa\rho \cos(\theta - \alpha) + \frac{1}{2}\pi i} \int_{-\infty}^{\sqrt{(2\kappa\rho) \cos \frac{1}{2}(\theta - \alpha)}} e^{-i\tau^2} d\tau + \frac{1}{\sqrt{\pi}} e^{i\kappa\rho \cos(\theta + \alpha) + \frac{1}{2}\pi i} \int_{-\infty}^{\sqrt{(2\kappa\rho) \cos \frac{1}{2}(\theta + \alpha)}} e^{-i\tau^2} d\tau,$$

and $l = \lambda \cos \alpha$, $n = \lambda \sin \alpha$, $\lambda = +\sqrt{l^2 + n^2}$, $\kappa = k\lambda$.

This agrees with the formulae obtained by Mr Clemmow.

REFERENCES

- Baker, B. B. & Copson, E. T. 1950 *Huygens' principle*, 2nd ed. Oxford: Clarendon Press.
 Bouwkamp, C. J. 1946 *Physica*, **12**, 467–474.
 Bouwkamp, C. J. 1947 *Math. Rev.* **8**, 179–180.
 Clemmow, P. C. 1950 (In preparation.)
 Copson, E. T. 1946a *Quart. J. Math.* **17**, 19–34.
 Copson, E. T. 1946b *Proc. Roy. Soc. A*, **186**, 100–118.
 Fox, E. N. 1948 *Phil. Trans. A*, **241**, 71–103.
 Fox, E. N. 1949 *Phil. Trans. A*, **242**, 1–32.
 Miles, J. W. 1949 *J. Appl. Phys.* **20**, 760–771.
 Rayleigh, Lord 1897 *Phil. Mag.* **43**, 259–272.
 Rayleigh, Lord 1913 *Proc. Roy. Soc. A*, **89**, 194–219.

The structure of linear relativistic wave equations. I

BY K. J. LE COUTEUR

Department of Theoretical Physics, University of Liverpool

(Communicated by P. A. M. Dirac, F.R.S.—Received 22 February 1950)

The structure of linear relativistic wave equations of the form $(i\alpha_\mu \partial^\mu - \chi) \psi = 0$ is discussed. In general, such equations describe particles with a spectrum of mass and spin values. It is proved that the physical requirement that all particle states can be unambiguously specified by the momentum, mass, spin and charge, leads to a complete determination of the eigenvalues of the total spin. These must form an arithmetical progression $S_H, S_H - 1, S_H - 2, \dots$, terminating at 0, $\frac{1}{2}$ or 1.

A coupling diagram is associated with every wave equation and necessary restrictions on the shape of the diagram are worked out. There results a useful reduction in the number of mathematical possibilities.

1. INTRODUCTION

Generalizations of Dirac's wave equation for the electron have been formulated by many authors. In this paper the mathematical possibilities are classified and discussed in a general way.

Relativistic notation is used with co-ordinates x_μ ($x_0 = ct$) and with the metric tensor $g_{\mu\nu}$ in the form

$$g_{00} = -g_{11} = -g_{22} = -g_{33} = 1.$$

Any set of relativistically invariant linear partial differential equations of finite order, with constant coefficients, can be written as a set of simultaneous linear partial differential equations of first order for a set of z quantities $\psi_s(x)$, $s = 1, 2, \dots, z$, which may be regarded as components of a vector $\psi(x)$ in some space of z dimensions. In matrix notation the equation is then

$$(B^\mu \partial_\mu - C) \psi(x) = 0. \quad (1.1)$$

Under the further assumption (α), that the B^μ and C are square matrices of order z and that C is non-singular, this equation can be written as

$$(\alpha^\mu p_\mu - 1\chi) \psi(x) = 0, \quad (1.2)$$

where χ is a non-vanishing constant. 1 is the unit matrix of order z ; in subsequent equations such factors will not be written explicitly. In the usual notation $p_\mu = i\partial_\mu$ or $-i\partial_\mu$, according as it operates to the right or to the left; where necessary an arrow indicates the direction of operation.

Dirac's equation and Kemmer's (1939) equation correspond to particular choices of the matrices α^μ in (1.2).

The restriction α , introduced above, is not satisfied in the systems of equations for particles of spin higher than 1 proposed by Dirac (1936) and Fierz & Pauli (1939). It has been shown by Bhabha (1945) that their equations are equivalent to (i) a set of equations of the type (1.2), and (ii) additional equations, not contained in the set (i), which impose relationships between the different components of $\psi(x)$ at each instant of time. Such relationships are known as 'subsidiary conditions', and they very seriously complicate the theory of the interaction between the ψ field and other fields. This aspect of the theory was considered by Fierz & Pauli (1939), who found it necessary to introduce auxiliary field variables in order to formulate the theory in a consistent way. Because of these difficulties some authors (Bhabha 1945; Kramers, Belinfante & Lubanski 1941; de Broglie 1943; Petiau 1946; Harish-Chandra 1947) have studied equations like (1.2) without imposing any additional subsidiary conditions. The field equations can then be derived as Lagrangian equations, and it is a straightforward matter to introduce interaction with other fields, usually by addition of the corresponding Lagrangians.

The main properties of relativistic wave equations of the form (1.2) can be elucidated, without detailed calculation, by the systematic use of simple algebraical and group theoretical arguments. The analysis shows that the physical requirements impose some necessary restrictions on the structure of the equation, and these enormously reduce the number of mathematical possibilities and so clarify the problem.

2. RELATIVISTIC INVARIANCE

Equation (1.2) is relativistically invariant if, with the Lorentz transformation

$$x' = ux, \quad p' = up \quad \text{or} \quad x'_\mu = u^\nu_\mu x_\nu, \quad p'_\mu = u^\nu_\mu p_\nu, \quad (2.1)$$

one can associate a linear transformation of the z components of ψ ,

$$\psi'(x') = U\psi(x), \quad (2.2)$$

such that

$$\alpha'_\mu = u^\nu_\mu \alpha_\nu = U^{-1} \alpha_\mu U. \quad (2.3)$$

Then any product of α matrices transforms as

$$\alpha'_\mu \alpha'_\nu \alpha'_\rho \dots = u^\theta_\mu u^\phi_\nu u^\psi_\rho \alpha_\theta \alpha_\phi \alpha_\psi \dots = U^{-1} \alpha_\mu \alpha_\nu \alpha_\rho \dots U. \quad (2.4)$$

This procedure is consistent only if the correspondence $u \rightarrow U$, defined by (2.1) and (2.2), obeys the law of composition $u \rightarrow U$, $t \rightarrow T$ implies $ut \rightarrow UT$. This condition

is satisfied (Wigner 1939), and so the Lorentz rotations (proper Lorentz transformations, of determinant +1, which do not reverse the time axis) can be built up from the infinitesimal transformations

$$x'_\mu = x_\mu + \epsilon_{\mu\nu} x^\nu, \quad U = 1 + \frac{1}{2} \epsilon^{\rho\sigma} I_{\rho\sigma}, \quad (2.5)$$

where $\epsilon_{\mu\nu}$ is small and $\epsilon_{\mu\nu} + \epsilon_{\nu\mu} = 0$. The $I_{\rho\sigma}$ are a set of six matrices of order z and $I_{\rho\sigma} + I_{\sigma\rho} = 0$. The condition $U\alpha'_\mu = \alpha_\mu U$ applied to (2.5) gives the basic commutation rule

$$(\alpha_\mu, I_{\rho\sigma}) = g_{\mu\rho} \alpha_\sigma - g_{\mu\sigma} \alpha_\rho, \quad (2.6)$$

where $(a, b) = ab - ba$. Similarly, the consistency condition $t^{-1}ut \rightarrow T^{-1}UT$, applied to (2.5), gives

$$T^{-1}I_{\mu\nu}T = t^\sigma_\mu t^\rho_\nu I_{\sigma\rho}, \quad (2.7)$$

which requires

$$(I_{\mu\nu}, I_{\rho\sigma}) = -g_{\mu\rho} I_{\nu\sigma} - g_{\nu\sigma} I_{\mu\rho} + g_{\mu\sigma} I_{\nu\rho} + g_{\nu\rho} I_{\mu\sigma}. \quad (2.8)$$

Note that (2.6) and (2.8) ensure the invariance of the theory under Lorentz rotations only. It is usually assumed, for reasons of symmetry (but also see § 7), that equation (1.2) takes the same form in left- and right-handed co-ordinate systems. If η_0 represents the U matrix corresponding to reversal of the three spatial axes, (2.3) and (2.7) require ($k, l = 1, 2, 3$),

$$\eta_0^{-1} \alpha_0 \eta_0 = \alpha_0, \quad \eta_0^{-1} \alpha_l \eta_0 = -\alpha_l, \quad (2.9)$$

$$\eta_0^{-1} I_{0l} \eta_0 = -I_{0l}, \quad \eta_0^{-1} I_{kl} \eta_0 = I_{kl}, \quad (2.10)$$

and these conditions, in conjunction with (2.5), ensure the invariance of (1.2) under improper Lorentz transformations. Since $(\eta_0)^2$ commutes with all the matrices α_μ and $I_{\mu\nu}$ it is permissible to take $(\eta_0)^2 = 1$.

3. REALITY CONDITIONS

The wave equation (1.2) and its complex conjugate may be derived from the scalar Lagrangian density

$$1 = \psi^\dagger(x) D(\alpha^\mu p_\mu - \chi) \psi(x), \quad (3.1)$$

provided that there exists a matrix D such that

$$(\alpha^\mu)^\dagger D = D \alpha^\mu \quad (3.2)$$

and

$$U^\dagger D U = D \quad \text{or} \quad U^\dagger D = D U^{-1},$$

where $(\alpha^\mu)^\dagger$ is the Hermitian conjugate of α^μ . By use of (2.5) for U one gets

$$I_{\rho\sigma}^\dagger D + D I_{\rho\sigma} = 0, \quad (3.3)$$

also

$$\eta_0^\dagger D = D \eta_0. \quad (3.4)$$

Wild (1947) has remarked that $D^\dagger$ and therefore $D + D^\dagger$ also satisfy these requirements; thus it can always be assumed that D is Hermitian.

4. COVARIANTS

The transformation (2.2) gives

$$\psi^\dagger D \rightarrow \psi'^\dagger D = \psi^\dagger U^\dagger D = \psi^\dagger D U^{-1}. \quad (4.1)$$

The corresponding transformation law of the covariant

$$s_{\mu\nu\rho} \dots = \psi^\dagger D \alpha_\mu \alpha_\nu \alpha_\rho \dots \psi \quad (4.2)$$

is

$$s_{\mu\nu\rho} \dots \rightarrow s'_{\mu\nu\rho} \dots = \psi'^\dagger D \alpha_\mu \alpha_\nu \alpha_\rho \dots \psi' = u_\mu^\theta u_\nu^\phi u_\rho^\psi \dots s_{\theta\phi\psi} \dots \quad (4.3)$$

by (4.1) and (2.4). This shows that $s_{\mu\nu\rho} \dots$ is a tensor. Special cases are the charge-current four-vector

$$s_\mu = \psi^\dagger D \alpha_\mu \psi, \quad (4.4)$$

and the energy momentum tensor

$$t_{\mu\nu} = \frac{1}{2} \{ \psi^\dagger D \cdot i \alpha_\mu \overrightarrow{\partial}_\nu \psi - \psi^\dagger D \cdot i \alpha_\mu \overleftarrow{\partial}_\nu \cdot \psi \}. \quad (4.5)$$

It is an important problem to determine the set of all linearly independent covariants associated with the wave equation (1.2). They are entirely determined by the number of components and transformation properties of ψ ; this is equivalent to saying that they are determined by the choice of the transforming matrices $I_{\mu\nu}$. For the covariant $\psi^\dagger D A \psi$, where A is any fixed matrix of order z , is just a linear combination of products $(\psi^\dagger D)_r \psi_s$, and so, in any Lorentz transformation (2.2), the set of all such covariants is transformed (Weyl 1931) by $U^{-1} \times U$, where the $\times$ denotes the Kronecker product.

It is usually assumed that the α -matrices and their multiple products generate the complete algebra of matrices of order z . Then the enumeration of all linearly independent covariants is equivalent to the enumeration of all linearly independent multiple products of the α -matrices, which is an essential step in the construction of the α -algebra.

5. THE LORENTZ GROUP IN FOUR DIMENSIONS

The spinor representations have been treated by Van der Waerden (1932) and Cartan (1938). This section contains a few simple definitions, and the formulae for the reduction of the direct product of two representations of the group are derived by a simplified method, analogous to the treatment of three-dimensional angular momenta (Dirac 1935, § 44).

The commutation rules (2.8) show that

$$\left. \begin{aligned} \frac{1}{2}(iI_{23} + I_{01}) &= P_1, & \frac{1}{2}(iI_{31} + I_{02}) &= P_2, & \frac{1}{2}(iI_{12} + I_{03}) &= P_3 \\ \frac{1}{2}(iI_{23} - I_{01}) &= Q_1, & \frac{1}{2}(iI_{31} - I_{02}) &= Q_2, & \frac{1}{2}(iI_{12} - I_{03}) &= Q_3 \end{aligned} \right\} \quad (5.1)$$

and

form two commuting sets of operators obeying the commutation rules,

$$(P_1, P_2) = iP_3, \quad (Q_1, Q_2) = iQ_3, \quad (5.2)$$

for angular momenta in three dimensions. Their magnitudes p and q are defined by

$$P_1^2 + P_2^2 + P_3^2 = p(p+1), \quad Q_1^2 + Q_2^2 + Q_3^2 = q(q+1), \quad (5.3)$$

and can also be derived from the scalar $I_{\mu\nu} I^{\mu\nu}$ and pseudo-scalar $\epsilon^{\mu\nu\sigma\rho} I_{\mu\nu} I_{\sigma\rho}$ by the equations

$$\left. \begin{aligned} \frac{1}{2}(I_{01}^2 + I_{02}^2 + I_{03}^2 - I_{23}^2 - I_{31}^2 - I_{12}^2) &= p(p+1) + q(q+1), \\ I_{23} I_{01} + I_{31} I_{02} + I_{12} I_{03} &= \frac{1}{i}(p(p+1) - q(q+1)). \end{aligned} \right\} \quad (5.4)$$

p and q commute with each other and with all the $I_{\mu\nu}$. Any finite irreducible representation of the group of Lorentz rotations can be specified by the eigenvalues, necessarily integral or half-odd integral, of p and q and is denoted by (p, q) . The order of this representation is $(2p+1)(2q+1)$.

To derive representations of the full Lorentz group it is necessary to include the reflexions. Equations (5.4) show that reflexion about the three spatial axes interchanges p and q so that, if $p \neq q$, the two representations (p, q) and (q, p) must occur together in a single irreducible representation of the Lorentz group with spatial reflexions, denoted by $D(p, q)$. Thus in $D(p, q)$ the representation of $I_{\mu\nu}$ is equivalent to the direct sum

$$I_{\mu\nu} = \begin{pmatrix} I_{\mu\nu}(p, q) & 0 \\ 0 & I_{\mu\nu}(q, p) \end{pmatrix}, \quad (5.5)$$

and then the reflexion operator η_0 , defined by (2.10), can be represented by

$$\eta_0(p, q) = a \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}. \quad (5.6)$$

The condition $\eta_0^2 = 1$ gives $a = +1$ or $a = -1$; these two representations for η_0 are equivalent, so that one can take $a = 1$. The reduction (5.5) of $I_{\mu\nu}$ implies a corresponding reduction of the U matrices built up as in (2.5); that is, for Lorentz rotations

$$D(p, q) = (p, q) \dot{+} (q, p). \quad (5.7)$$

The representation (p, p) , however, gives rise to two different representations $D^+(p, p)$ and $D^-(p, p)$ of the Lorentz group with reflexions. Both are of order $(2p+1)^2$ and for Lorentz rotations

$$D^+(p, p) = (p, p) = D^-(p, p). \quad (5.8)$$

D^+ and D^- correspond to different choices of the sign of the reflexion operator, so that

$$\eta_0^+(p, p) = -\eta_0^-(p, p). \quad (5.9)$$

Equation (5.19) gives $|\text{spur } \eta_0^\pm(p, p)| = (2p+1)$; since this is non-zero, η_0^+ and η_0^- are inequivalent representations of the spatial reflexion.

The transformation (2.2) must be generated by the direct sum of a number of irreducible representations of the Lorentz group. In the representation with p and q diagonal, the matrices $I_{\mu\nu}$ and the vector ψ take the form

$$I_{\mu\nu} = \begin{pmatrix} I_{\mu\nu}(a) & . & . & . \\ . & I_{\mu\nu}(b) & . & . \\ . & . & I_{\mu\nu}(c) & . \\ . & . & . & . \end{pmatrix}, \quad \psi = \begin{pmatrix} \psi(a) \\ \psi(b) \\ \psi(c) \\ . \end{pmatrix} \quad (5.10)$$

and each $I_{\mu\nu}(a)$ is an irreducible representation (p_a, q_a) of the Lorentz rotations.

Any Lorentz rotation induces simultaneous transformations

$$\psi'(a) = U(a)\psi(a), \quad \psi'(b) = U(b)\psi(b) \quad (5.11)$$

of $\psi(a)$ and $\psi(b)$, and the set of all products $\psi_r(a)\psi_s(b)$ is transformed by $U(a) \times U(b)$. If $U(a)$ and $U(b)$ are in the infinitesimal form (2.5), we obtain

$$\begin{aligned} U(a) \times U(b) &= 1 + \frac{1}{2}1(a) \times \epsilon^{\rho\sigma} I_{\rho\sigma}(b) + \frac{1}{2}\epsilon^{\sigma\rho} I_{\rho\sigma}(a) \times 1b \\ &= 1 + \frac{1}{2}\epsilon^{\rho\sigma} I_{\rho\sigma}. \end{aligned} \quad (5.12)$$

This shows that the eigenvalues of $P_1 = \frac{1}{2}(iI_{23} + I_{01})$ are obtained by taking all possible combinations of an eigenvalue of $P_1(a)$ with an eigenvalue of $P_1(b)$; therefore (Dirac 1935, § 44) P may be represented as the direct sum over the representations $P = p_a + p_b$ to $P = |p_a - p_b|$ of the three-dimensional rotation group. Similar results hold independently for Q . The generating elements $I_{\mu\nu}$ of the product transformation $U(a) \times U(b)$ thus consist of the direct sum of a number of irreducible sets of $I_{\mu\nu}$, and this is written symbolically as

$$(p_a, q_a) \times (p_b, q_b) = \Sigma(p, q), \quad (5.13)$$

with $p_a + p_b \geq p \geq |p_a - p_b|$ and $q_a + q_b \geq q \geq |q_a - q_b|$.

This equation, of course, remains valid if the labels a, b refer to the same representation.

Equation (5.13) is easily extended to the reduction of the direct product of two representations of the Lorentz group with reflexions. Suppose $p_a \neq q_a, p_b \neq q_b$, then

$$(p_a, q_a) \dot{+} (q_a, p_a) \times (p_b, q_b) \dot{+} (q_b, p_b) = \sum^1(p, q) + (q, p) \dot{+} \sum^2(p, q) + (q, p) \quad (5.14)$$

over (1) $p_a + p_b \geq p \geq |p_a - p_b|, \quad q_a + q_b \geq q \geq |q_a - q_b|$

and (2) $p_a + q_b \geq p \geq |p_a - q_b|, \quad q_a + p_b \geq q \geq |q_a - p_b|$.

In the product representation, the reflexion operator is given by

$$\eta_0 = \eta_0(a) \times \eta_0(b),$$

and by changing the sign of $\eta_0(a)$ one gets an equivalent representation of the direct product and so of η_0 . But this changes the sign of η_0 , therefore in the decomposition of η_0 the representations $\eta_0^+(p, p)$ and $\eta_0^-(p, p)$, if they occur, appear together in pairs. Therefore

$$D(p_a, q_a) \times D(p_b, q_b) = \sum^1 D(p, q) \dot{+} \sum^2 D(p, q), \quad (5.15)$$

and whenever a pair of (p, p) representations occur in Σ^1 or Σ^2 it must be written as $D^+(p, p) \dot{+} D^-(p, p)$. Similarly, with the same convention,

$$D(p_a, q_a) \times D(p_b, p_b) = \Sigma D(p, q), \quad (5.16)$$

with $p_a + p_b \geq p \geq |p_a - p_b|$ and $q_a + p_b \geq q \geq |q_a - p_b|$.

Similarly, if x and y stand for $+$ and $-$,

$$D^x(p_a, p_a) \times D^y(p_b, p_b) = \Sigma D^{xy}(p, q), \quad (5.17)$$

with

$$p_a + p_b \geq p \geq q \geq |p_a - p_b|,$$

and the sign xy on the right-hand side applies only if $p = q$. The association of signs in (5.17) constitutes a convention for identifying η_0^+ and η_0^- in (5.9), for (5.17) gives

$$\text{spur } \eta_0^+(p_a, p_a) \times \text{spur } \eta_0^+(p_b, p_b) = \Sigma \text{spur } \eta_0^+(p, p). \quad (5.18)$$

The ψ which carries the representation $D^+(\frac{1}{2}, \frac{1}{2})$ is a four-vector; under spatial reflexion three components change sign so that $\text{spur } \eta_0^+(\frac{1}{2}, \frac{1}{2}) = -2$, therefore the solution of (5.18) is

$$\text{spur } \eta_0^+(p, p) = (-)^{2\nu} (2p + 1). \quad (5.19)$$

These reduction formulae (5.13 to 5.19) enable one to evaluate the direct product $U \times U'$, where U and U' are any two matrix transformations (2.2) associated with the same Lorentz transformation (2.1). They remain valid if $U(a)$ or $U(b)$ is replaced by the reciprocal transformation; for this is merely equivalent to changing the sign of $I_{\mu\nu}(a)$ or $I_{\mu\nu}(b)$, which does not affect the combination of the two sets of spin operators. In particular, the formulae can be applied to the direct product $U^{-1} \times U$ which was shown in § 4 to determine all the covariants associated with equation (1.2) and to characterize the algebra of the α -matrices.

Example. Dirac's equation

ψ transforms according to the representation $D(\frac{1}{2}, 0)$ and

$$D(\frac{1}{2}, 0) \times D(\frac{1}{2}, 0) = D^+(0, 0) \dot{+} D^-(0, 0) \dot{+} D^+(\frac{1}{2}, \frac{1}{2}) \dot{+} D^-(\frac{1}{2}, \frac{1}{2}) \dot{+} D(1, 0),$$

so that the covariants consist of 1 scalar, 1 pseudoscalar, 1 vector, 1 pseudovector and 1 skew-tensor of second rank. It follows that the linearly independent elements of the α -algebra are

$$1, \quad \omega = \epsilon^{\mu\nu\sigma\rho} \alpha_\mu \alpha_\nu \alpha_\sigma \alpha_\rho, \quad \alpha_\mu, \quad \omega \alpha_\mu, \quad (\alpha_\mu, \alpha_\nu),$$

and since the symmetric second rank tensor $\alpha_\mu \alpha_\nu + \alpha_\nu \alpha_\mu$ does not appear, it must be fully reducible, that is,

$$\alpha_\mu \alpha_\nu + \alpha_\nu \alpha_\mu = kg_{\mu\nu},$$

which are well-known commutation rules of Dirac's theory.

The same method can be used to list the covariants and to derive the α -algebra in the scalar and vector meson theories.

6. MASS AND SPIN EIGENVALUES

It is known (Harish-Chandra 1947) that the characteristic equation of α_0 must be of the form

$$\alpha_0(\alpha_0^2 - j_1^2)(\alpha_0^2 - j_2^2) \dots = 0, \quad (6.1)$$

with real coefficients j . The factor α_0 occurs only if α_0 is singular. Wild (1947) has shown that if the canonical form of α_0 is not diagonal the equation (1.2) contains implicit subsidiary conditions which lead to difficulties. Therefore only α -matrices reducible to diagonal form are considered and then no factor in (6.1) need appear more than once.

The form of (6.1) follows from the invariance of (1.2) under the total reflexion represented by a U -matrix η such that

$$\alpha'_\mu = -\alpha_\mu = \eta^{-1} \alpha_\mu \eta, \quad I'_{\mu\nu} = I_{\mu\nu} = \eta^{-1} I_{\mu\nu} \eta. \quad (6.2)$$

It will be shown in § 8 that η always exists, even though (1.2) is not necessarily invariant against spatial reflexions.

α_0 commutes with I_{23} , I_{31} , I_{12} , and so with each of the irreducible representations D_s of the three-dimensional rotation group, which are labelled by the eigenvalues s of S , defined by

$$S(S+1) = -(I_{23}^2 + I_{31}^2 + I_{12}^2). \quad (6.3)$$

With each eigenvalue j of α_0 is associated a set of eigenvectors $\psi(r, j)$, where the parameter r defines the simultaneous eigenvalues of S and of any other matrices, such as I_{23} , which can be combined with α_0 and S to form a complete commuting set. The transformation (6.2) shows that, if $j \neq 0$, the eigenvalue $-j$ must appear with another equivalent set of representations D_s . It is for this reason that every particle must have an anti-particle (Pauli 1940). The transformation also shows that if some D_s occurs an odd number of times, the matrices α_μ must be singular.

The spatially constant solutions of (1.2) are

$$\psi = \psi(r, j) e^{-i(\chi/j)x^0}, \quad (6.4)$$

where j is any non-zero eigenvalue of α_0 . In a quantum-mechanical interpretation of (1.2) this represents the wave function of a free particle with zero momentum and energy (χ/j) . In this state the particle has mass $|\chi/j|$ and spin angular momentum $\hbar s$, and so to each wave equation (2.2) corresponds a spectrum of mass-spin eigenvalues.

The set of spin representations D_s associated with (1.2) is defined by the reduction (5.10) of the $I_{\mu\nu}$. For iI_{23} , iI_{31} , iI_{12} is the resultant of P and Q , defined by (5.1), and so the representation (p, q) contains the representations

$$\Sigma D_s, \quad p+q \geq s \geq |p-q| \quad (6.5)$$

of the three-dimensional rotation group. But only the representations D_s which occur with non-zero eigenvalues j of α_0 can be interpreted, through (6.4), as spin angular momenta. The exceptional eigenvalues of S which occur only with vanishing eigenvalues of α_0 have no such interpretation; however, this possibility is confined to particular values of s which are given in § 10.

7. INVARIANCE UNDER SPATIAL REFLEXIONS

In § 2, it was seen that equations invariant under Lorentz rotations are not necessarily invariant under spatial reflexions, which must be introduced through an additional assumption.

In §§ 3 and 4, in order to define a real energy momentum tensor and charge current four-vector, it was necessary to introduce the matrix D , defined by (3.2) and (3.3). Now, § 5 shows that it is always possible to choose the representation so that iI_{kl} and I_{0l} are Hermitian matrices; then D must commute with I_{kl} and anti-commute with I_{0l} . The matrix elements of D are therefore of the form $(p, q | D | q, p)$,

connecting different representations of the Lorentz rotations. If D is to exist, the representations (p, q) associated, as in (5.10), with the wave equation must appear in pairs, making up complete representations $D(p, q)$ of the full Lorentz group. Thus a reflexion operator η_0 , satisfying the commutation rules (2.10), must exist, and it is natural to require that equations (2.9) are also satisfied, so making the wave equation (1.2) invariant under all Lorentz transformations.

In a representation such as that of § 5 with iI_{kl} , I_{0l} and η_0 Hermitian, $D\eta_0^{-1} = K$ commutes with all the $I_{\mu\nu}$ and with η_0 . Thus (Wild 1947)

$$D = K\eta_0, \quad (7.1)$$

where K is a scalar, represented in each $D(p, q)$ by a multiple $K(p, q)$ of the unit matrix. Wild has shown that, without loss of generality, one may assume

$$K(p, q) = \pm 1.$$

The existence of D ensures that α_0 is equivalent to its Hermitian conjugate, but it does not follow that the eigenvalues of α_0 are all real, as is required for the physical interpretation of § 6. However, in the above representation, (7.1) and (3.2) give

$$\alpha_0^\dagger K = K\alpha_0, \quad (7.2)$$

and, if only the simple case $K = 1$ is considered, α_0 is Hermitian and so automatically has real eigenvalues. This restriction on the choice of K is convenient, but not always necessary.

8. THE α -MATRICES

The results of § 5 can be used to determine the form of the α -matrices in the representation (5.10) with p and q diagonal. For since $\psi^\dagger D\alpha_\mu \psi$ is a four-vector, § 4 shows that α_μ can have matrix elements $(a | \alpha_\mu | b)$ connecting representations a and b only if the representation $(\frac{1}{2}, \frac{1}{2})$ occurs in the reduction (5.13) of the direct product. This argument, due to Garding, gives a condition

$$p_a - p_b = \pm \frac{1}{2}, \quad q_a - q_b = \pm \frac{1}{2}, \quad (8.1)$$

and when it is satisfied $(\frac{1}{2}, \frac{1}{2})$ occurs only once in the reduction, so that the part of α_μ which connects a and b is unique apart from a multiplying constant, C_{ab} , and may be written as

$$(a | \alpha^\mu | b) = C_{ab}(a | a^\mu | b). \quad (8.2)$$

Therefore, with some convention about the form of the basic coupling matrices $(a | a^\mu | b)$, a relativistic wave equation can be specified by the set of its coupling coefficients.

The form (8.1) of the α^μ shows that the reflexion matrix η of (6.2) always exists and has matrix elements

$$\eta_a = +1 \text{ for } p_a \text{ integral, } \eta_a = -1 \text{ for } p_a \text{ half-odd integral.} \quad (8.3)$$

It is noteworthy that these results assume only invariance against Lorentz rotations.

The rules (5.13) can be used in a similar way to study the matrix elements of simple combinations of the α -matrices, for example, of the brackets (α_μ, α_ν) .

It is convenient to arrange the points (p, q) in a rectangular lattice and then the individual couplings $(p, q) \leftrightarrow (p \pm \frac{1}{2}, q \pm \frac{1}{2})$ appear as diagonals joining opposite

corners of a unit cell (see figure 1). Thus a diagram is associated with every relativistic wave equation, and these are particularly useful if no representation (p, q) occurs more than once with the wave equation. If the wave equation is invariant under spatial reflexion, the diagram must be unchanged by the reflexion $(p, q) \leftrightarrow (q, p)$ about the symmetry axis $p = q$.

It is usual (Bhabha 1945) to assume that the wave equation (1.2) cannot be broken up into a number of independent equations for different subvectors of ψ ; this is permissible since any equation can be treated as a combination of such equations. This requires that the diagram consists either of a single connected piece or else of two separate pieces which go over into each other by reflexion.

These considerations restrict the possible sets of spin eigenvalues. For (8.1) shows that, in a connected diagram, either each $p_a + q_a$ is integral or each $p_a + q_a$ is half-odd integral so that (6.5) the eigenvalues s of S are either all integral or all half-odd integral. However, apart from this sharp separation into two classes, it is not in general possible to label a wave equation by any particular spin value. The greatest eigenvalue $s_{\max.}$ of S is given by the greatest value of $p + q$ and the least eigenvalue $s_{\min.}$ of S by the least value of $|p - q|$ which occurs in (5.10); the other eigenvalues of S form a progression

$$s_{\max.}, (s_{\max.} - 1), \dots, s_{\min.}, \quad (8.4)$$

from which no member may be missing. It will appear below (§ 10) that the value of $s_{\min.}$ is fixed by physical considerations.

Given any definite representation of the $I_{\mu\nu}$, the coupling matrices can be worked out from (2.6). It is sufficient to determine α_0 , for then α_l is given by

$$\alpha_l = (\alpha_0, I_{0l}). \quad (8.5)$$

The form of α_0 or a_0 can be worked out from the conditions

$$(a_0, I_{kl}) = 0, \quad ((a_0, I_{0l}), I_{0l}) = a_0. \quad (8.6)$$

In a spinor representation Bhabha (1945) expressed the basic coupling matrices in terms of Dirac's (1936) spinor matrices $u^\alpha(k)$ and $v^\alpha(k)$. For some purposes, especially for physical interpretation in terms of eigenvalues of mass and spin, other representations are more convenient and will be described in a second paper.

The representation can be constructed so that the I_{0l} and iI_{kl} are Hermitian matrices, and all the $I_{\mu\nu}$ are either purely real or purely imaginary, it is easy to see that this is consistent with the commutation rules (2.8). Then (8.6) shows that $\alpha_0^\dagger$ and $\bar{a}_0$ also satisfy the conditions on a_0 ; therefore, it is possible to choose the basic coupling matrices so that

$$\left. \begin{array}{l} \text{(i) } (a | a^0 | b) \text{ is purely real,} \\ \text{(ii) } (a | a^0 | b) \text{ is the Hermitian conjugate of } (b | a^0 | a). \end{array} \right\} \quad (8.7)$$

The condition $\alpha_0^\dagger = D\alpha_0 D^{-1} = K\alpha_0 K^{-1}$

by (7.2), then gives $\bar{C}_{ba} = K_a C_{ab} K_b = \pm C_{ab}, \quad (8.8)$

so that each diagonal coupling $a \leftrightarrow b$ involves essentially only one coupling coefficient.

9. AN ADDITIONAL PHYSICAL RESTRICTION

The foregoing scheme is very general; fortunately, it can be made more definite by considering the requirements of the physical interpretation.

No meaning can be given to the transition of a particle between two states which are observationally indistinguishable, therefore the formulations of particle wave equations should be restricted by the requirement (Le Couteur 1949, § 13) that all the particle states involved in the equation can be unambiguously specified by dynamical variables of real physical significance. In § 6 it was seen that for a free particle, considered in the frame of reference of zero spatial momentum, the natural variables are the mass, the sign of the charge and the spin variables. There is no ambiguity if and only if each combination of these quantities occurs only once; according to § 6 this requires that:

The representations D_s which occur in conjunction with the same, non-zero, eigenvalue j of α_0 are all inequivalent. The sign of j is taken into account as it represents the sign of the charge.

If this restriction is accepted, the task of constructing an irreducible relativistic wave equation with an assigned mass-spin spectrum is a quite definite one; otherwise the problem would be completely indeterminate.

10. THE COUPLING DIAGRAMS

The requirements of § 9 leads to stringent restrictions on the form of the coupling diagram.

The condition on α_0 requires that the diagram consists of a single connected piece rather than of two pieces which are exchanged by reflexion. There is, further, the stronger requirement that, in each of the different subdiagrams formed by considering only the representations (p, q) which contain a given value of s , every representation (p, q) is either

(a) connected to the corresponding (q, p) by a chain of diagonal couplings lying entirely inside the subdiagram, or

(b) (p, q) is not connected to any other representation by couplings lying entirely inside the subdiagram.

For if, for some value of s , some of the representations (p, q) in the subdiagram do not satisfy condition (a), the matrix $(s | \alpha^0 | s)$ contains at least one pair of equivalent unconnected submatrices which are transformed into one another by the reflexion operator η_0 . Each pair of submatrices gives rise to pairs of equal eigenvalues which occur with equivalent representations D_s . The restriction on α_0 is violated unless all these repeated eigenvalues vanish. In this case, since it is assumed that the canonical form of α_0 is diagonal, each pair of equivalent submatrices must vanish completely, and so the representations (p, q) which do not satisfy condition (a) have to satisfy condition (b).

Figure 1 illustrates some allowed and some forbidden diagrams. Diagram 1b is forbidden as, for $s = \frac{3}{2}$, it violates condition (a).

Since the diagram forms a single connected piece, it must cross the axis of symmetry $p = q$. It follows that the least eigenvalue $s_{\min.}$ of S , which was introduced in (8.4), is $s_{\min.} = 0$ or $s_{\min.} = \frac{1}{2}$, according as the spin is integral or half-odd integral.

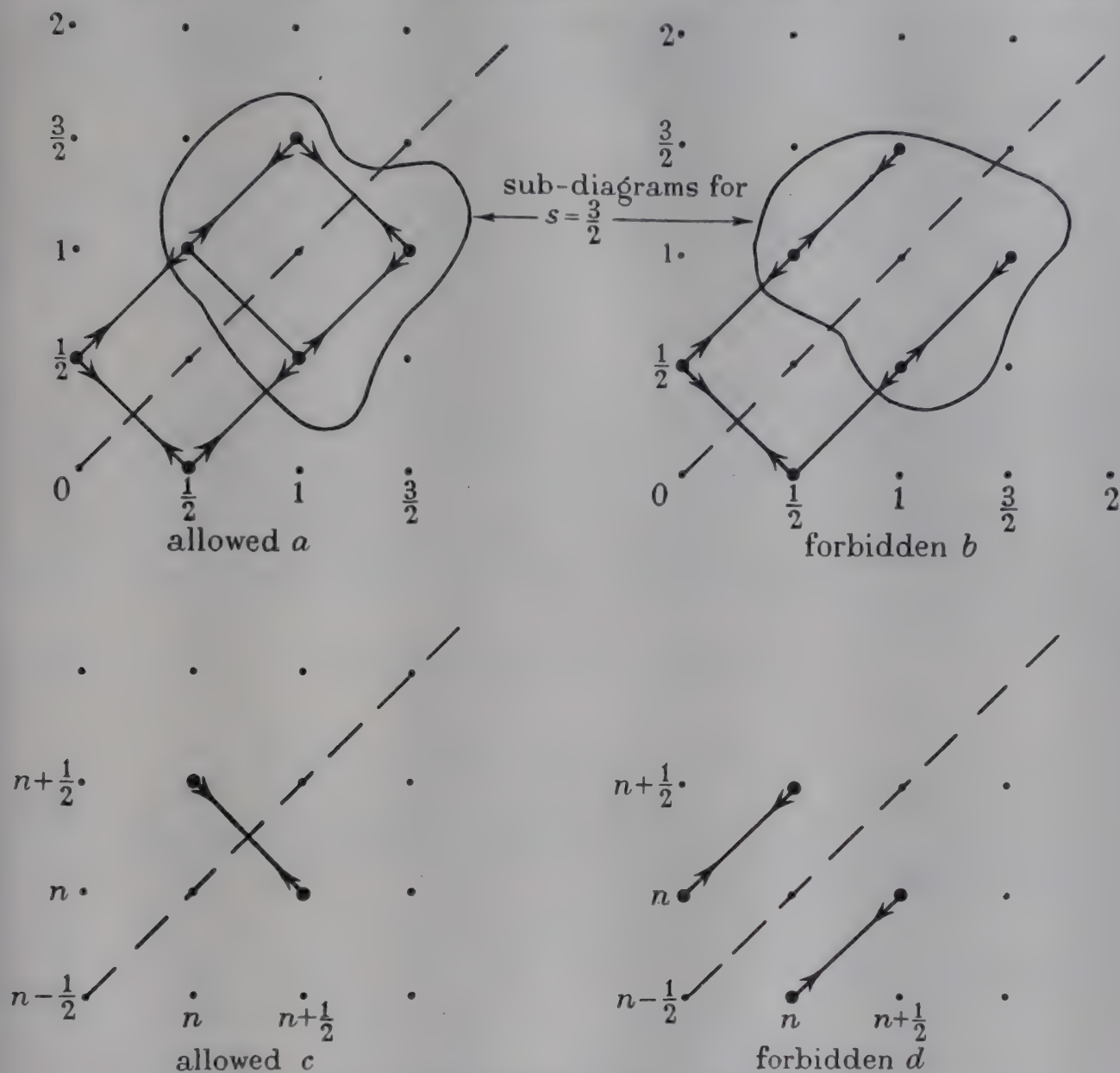


FIGURE 1. Allowed (a, c) and forbidden (b, d) forms of the coupling diagram.

It is now possible to identify the exceptional eigenvalues of S , that is (§ 6), those which occur only with zero eigenvalues of α_0 . Let S_H be the greatest value of s which occurs with a non-zero eigenvalue of α_0 . Thus S_H represents the highest possible value for the spin angular momentum of the particle described by the wave equation. Then, if $s_H < s_{\max.}$, all the submatrices $(s | \alpha^0 | s)$ of α^0 with $s_H < s \leq s_{\max.}$ must vanish completely. Since the representations containing $s_{\max.}$ must be connected to those containing the lower values of s by diagonal links like (8.1), the only possibilities are

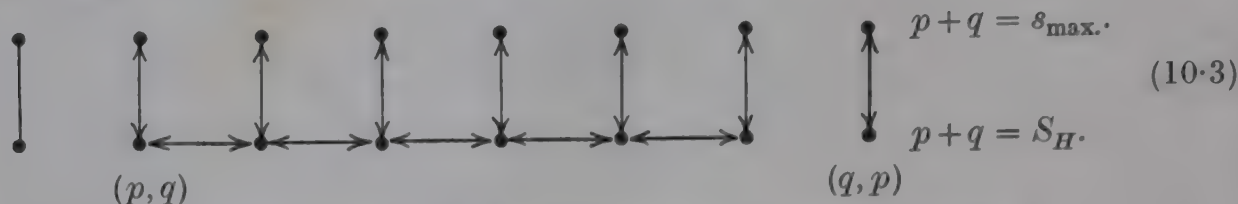
$$s_{\max.} = S_H \quad \text{or} \quad s_{\max.} = S_H + 1. \quad (10.1)$$

For $S_H = s_{\max.}$, the S_H subdiagram is of the form

$$(p' q') \quad (p, q) \quad (p - \frac{1}{2}, q + \frac{1}{2}) \quad (q, p) \quad (q', p')$$

with $p+q = S_H, \quad p'+q' = S_H, \quad p' > p > q. \quad (10.2)$

Condition (a) above requires that the chain $(p, q) \leftrightarrow (q, p)$ is complete. The uncoupled representation like (p', q') appear only if α_0 has several zero eigenvalues associated with $s = S_H$; in the full diagram they must be coupled to other representations by links like $(p', q') \leftrightarrow (p' - \frac{1}{2}, q' - \frac{1}{2})$. For $s_{\max.} = S_H + 1$, the S_H subdiagram contains additional couplings which link the representations with $p+q = s_{\max.}$ to those with $p+q = S_H$.



Condition (a) again requires that the chain $(p, q) \leftrightarrow (q, p)$ is complete.

The chains contained in (10.2) or (10.3) contribute non-zero couplings to every subdiagram with $s \leq S_H$ except, in the case of integral spins, the subdiagram $s = 0$. Therefore $(s | \alpha^0 | s)$ cannot vanish completely except perhaps for $s = 0$ and $s = s_{\max.}$. Equation (8.4) can now be replaced by the more physical statement that the eigenvalues of the spin angular momentum given by the wave equation (1.2) form a complete progression

$$S_H, \quad (S_H - 1), \quad \dots, \quad \begin{pmatrix} 0 \\ \frac{1}{2} \\ 1 \end{pmatrix}, \quad (10.4)$$

ending with $0, \frac{1}{2}, 1$. No member of this sequence may be omitted.

The vector and scalar meson theories provide examples of the two possible ways in which eigenvalues of S can fail to be realized as spin angular momenta.

Vector theory: $(1, 0) \leftrightarrow (\frac{1}{2}, \frac{1}{2}) \leftrightarrow (0, 1)$.

$s = 1$ is realized, $s = 0$ is not.

Scalar theory: $(0, 0) \leftrightarrow (\frac{1}{2}, \frac{1}{2})$.

$s = 0$ is realized, $s = 1 = s_{\max.}$ is not.

These conditions on the form of the diagram are necessary, but, of course, not sufficient. For example, in Bhabha's (1945) theory, the representations (p, q) associated with his wave equation $R_5(n, m)$ form the array

$$\left. \begin{array}{ccc} \left(\frac{n+m}{2}, \frac{n-m}{2} \right) & \dots & \left(\frac{n-m}{2}, \frac{n+m}{2} \right) \\ \left(\frac{n+m-1}{2}, \frac{n-m-1}{2} \right) & \dots & \left(\frac{n-m-1}{2}, \frac{n+m-1}{2} \right) \\ \dots & & \dots \\ (m, 0) & \dots & (0, m) \end{array} \right\} \quad (10.5)$$

which satisfies conditions (a). However, it has been shown (Le Couteur 1949, appendix) that the physical requirement of § 9 is not satisfied unless

$$m = 0 \quad \text{or} \quad m = \frac{1}{2} \quad \text{or} \quad m = n. \quad (10.6)$$

11. NON-RELATIVISTIC APPROXIMATION AND MAGNETIC MOMENT

If the particle described by (1.2) has electric charge, it will also show a magnetic moment which can be worked out in terms of the α -matrices. The simplest form of interaction with the electromagnetic field is obtained by substituting

$$p^\mu = i\partial^\mu + eA^\mu \quad (11.1)$$

in equation (1.2) and (3.1). A^μ is a vector potential of the electromagnetic field and e is a constant which determines the charge of the particle in all its states. The equations then admit the gauge transformation

$$\psi \rightarrow \{\exp(i e \int A'_\mu dx^\mu)\} \psi, \quad A_\mu \rightarrow A_\mu + A'_\mu. \quad (11.2)$$

Let us work in a representation with α_0 diagonal, and use $j, j', j'', \dots$ to denote eigenvalues of α_0 . Then if j is a particular eigenvalue of α_0 , the wave equation (1.2) has solutions of the form

$$\psi = e^{-i(\chi/j)x^0} \phi$$

$$\text{if} \quad \left\{ \left(\frac{\alpha_0}{j} - 1 \right) \chi + p^0 \alpha_0 + p^k \alpha_k \right\} \phi = 0, \quad (11.3)$$

$$\text{which gives} \quad p^0 j \phi(j) + p^k \sum^{j'} (j | \alpha_k | j') \phi(j') = 0, \quad (11.4a)$$

$$\left(\frac{\chi}{j} - \frac{\chi}{j'} + p^0 \right) j' \phi(j') + p^k \sum^{j''} (j' | \alpha_k | j'') \phi(j'') = 0. \quad (11.4b)$$

It is now assumed that p^k and p^0 are small compared with $(\chi/j - \chi/j')$, the difference between the rest energies in the states of mass χ/j and χ/j' . This condition is fulfilled if (χ/j) is the lowest mass eigenvalue of the particle, since it cannot jump into a state of higher mass unless the external field A^μ contains energetic photons. The important component of ϕ is then $\phi(j)$, the others are of the first order of smallness and so (11.4b) gives

$$\phi(j') = \frac{p^k}{\chi} \frac{j}{j-j'} (j' | \alpha_k | j) \phi(j),$$

and then, by substitution into (11.4a),

$$\left\{ p^0 + \frac{p^k p^l}{\chi} \sum^{j'} (j | \alpha_k | j') \frac{1}{j-j'} (j' | \alpha_l | j) \right\} \phi(j) = 0. \quad (11.5)$$

Now, according to (8.5),

$$(j | \alpha_l | j') = (j-j') (j | I_{0l} | j'),$$

$$\text{so that} \quad \sum^{j'} (j | \alpha_2 | j') \frac{1}{j-j'} (j' | \alpha_3 | j) = \frac{1}{2} (j | I_{02} \alpha_3 - \alpha_2 I_{03} | j), \quad (11.6)$$

$$\begin{aligned} \text{and} \quad \sum^{j'} (j | \alpha_2 | j') \frac{1}{j-j'} (j' | \alpha_2 | j) &= \frac{1}{2} (j | I_{02} \alpha_2 - \alpha_2 I_{02} | j) \\ &= -\frac{1}{2} (j | \alpha_0 | j) = -\frac{1}{2} j. \end{aligned} \quad (11.7)$$

The wave equation (11.5) now becomes

$$\left\{ p^0 - \frac{j}{2\chi} \mathbf{p}^2 + \sum_{1,2,3} (p^2 p^3 - p^3 p^2) \frac{1}{2\chi} (j | I_{02} \alpha_3 - I_{03} \alpha_2 | j) \right\} \phi(j) = 0.$$

Now $(p^2 p^3 - p^3 p^2) = (i\partial^2 + eA^2, i\partial^3 + eA^3) = ieH^1,$

where H is the magnetic field strength. So, with $M = \chi/j$,

$$\left\{ p^0 - \frac{1}{2M} \mathbf{p}^2 + \Sigma H^1 \frac{e}{2M} \frac{i}{j} (j | I_{02} \alpha_3 - I_{03} \alpha_2 | j) \right\} \phi(j) = 0, \quad (11.8)$$

which is the Schrödinger equation for a particle of mass M and magnetic moment

$$\frac{e}{M} \frac{i}{2j} (j | I_{02} \alpha_3 - I_{03} \alpha_2 | j), \quad \text{etc.} \quad (11.9)$$

The calculation is similar to the corresponding calculations in the electron and meson theories and is Bhabha's (1945) theory. However, (11.9) shows that in general the magnetic moment in the ground state of lowest rest mass is not simply related to the mass in that state. Explicit calculations require a representation of the matrices and are given in a second paper. It is possible to set up a simple model which correctly represents the magnetic moment of the proton.

12. NOTE ON THE ELECTROMAGNETIC INTERACTION

Except in the simplest cases, the α -algebra is large enough to contain amongst the covariants (§ 4) several scalar elements E , apart from simple multiples of the unit matrix. This happens even for the simple scalar-meson equation. It is important to consider whether these can be used to generalize the electromagnetic interaction (11.1) to

$$p^\mu = i\partial^\mu + EA^\mu, \quad (12.1)$$

which seems to offer a possibility of using a single irreducible wave equation to describe particles having different charge as well as different mass states.

The Lagrangian (3.1) is real provided

$$D\alpha^\mu E = E^\dagger \alpha^{\mu\dagger} D = E^\dagger D\alpha^\mu. \quad (12.2)$$

The equations still admit the gauge transformation (11.2) with E in place of e and the Maxwell equations show that the charge-current four-vector must be

$$s_\mu = \frac{\partial L}{\partial A^\mu} = \psi^\dagger D\alpha_\mu E \psi. \quad (12.3)$$

Then the equations of motion for ψ and $\psi^\dagger$ give

$$\partial^\mu s_\mu - i\chi\psi^\dagger (DE - E^\dagger D)\psi = 0, \quad (12.4)$$

so that unless a subsidiary condition is imposed on ψ , the conservation of charge requires

$$DE = E^\dagger D. \quad (12.5)$$

This excludes all the more general solutions E of (12.2), for the equations together give

$$\alpha_\mu E = E \alpha_\mu,$$

which shows that, if the α_μ form an irreducible set of matrices, E must be a simple multiple e of the unit matrix. It is therefore not possible to use one irreducible wave equation to describe particles with different charge states.

13. DISCUSSION

In general, irreducible relativistic wave equations more complicated than those of Dirac and Kemmer describe particles which have a spectrum of mass and spin values. It is remarkable that the physical requirement that the states of a free particle can be fully described by the momentum, mass, spin and charge, without the addition of other variables, leads to a complete determination of the form (10.4) of the spin spectrum. The different mass states must be very similar to each other; for example, if in one state the particle interacts strongly with some external field, then in all states it can interact with this field. Thus this formalism might possibly be applied to describe a family of mesons, all of which interact strongly with nucleons. There is some slight experimental evidence that such a family of mesons does exist; also theoretically, it is very difficult to account for the nuclear forces by using only a single meson field.

Since the wave equation is supposed irreducible, the α -matrices and their products must have matrix elements connecting the states of different rest mass. Thus, in the presence of interaction with other fields, the states of higher rest mass are unstable and have short lifetimes for decay into states of lower rest mass. An example of decay by emission of a γ -ray has been worked out (Le Couteur 1949).

This mode of decay is not in itself an objection to the use of a single irreducible wave equation to describe a family of nuclear force mesons. Even if each member of such a family is represented by a quite independent wave equation, the heavy mesons can decay into lighter mesons by virtual nuclear processes. The lifetime is very short (Van Wyk 1949).

The determination of the form (10.4) of the spin spectrum is of some practical importance. For if one attempts to find an irreducible wave equation for a particle with several states of mass and spin, knowledge of the total number of states to be described at once sets an upper limit to the spin and so confines the coupling diagram to a small area in the corner of the (p, q) plane. For example, if only one state has to be described, the spin cannot be higher than one. Without this principle there would be many more mathematical possibilities, for the spin is not easily measured directly.

For some forms of the coupling diagram, the mass values are not all independent. Examples of 'non-trivial' relationships between the mass values, that is to say relationships independent of the choice of the coupling parameters, are worked out in a second paper.

Subsidiary conditions were introduced into the older theories of Dirac, Fierz and Pauli. They consist of additional equations which confine the free plane-wave solutions (6.4) of (1.2) to a single mass and spin value and make either the energy

or the charge positive definite. Such conditions are consistent with the equations of motion for free particles, since each mass-spin state is a stationary state. On the other hand, if interactions are introduced, they have matrix elements connecting states with different mass and spin values and so are inconsistent with the simplest form of the subsidiary conditions. Fierz & Pauli (1939) have shown that to obtain consistent equations in the presence of interaction with the electromagnetic field, it is necessary to reformulate the equations in a more complicated way. Quantization is then very difficult. Since these older theories were formulated, our views on the existence of elementary particles have changed so much that there is no longer any reason to complicate matters by adding subsidiary conditions to make the mass and spin unique. Therefore, the theory described in this paper contains no subsidiary conditions, and, as pointed out in the introduction, this simplifies the treatment of interaction with other fields. The quantization of the theory has already been described (Le Couteur 1949).

REFERENCES

- Bhabha, H. J. 1945 *Rev. Mod. Phys.* **17**, 200.
 Cartan, E. 1938 *Actualités sci. industr.* **643** and **701**.
 de Broglie, L. 1943 *Théorie générale des particules à spin*. Paris: Gauthier Villars.
 Dirac, P. A. M. 1935 *Quantum mechanics*. Oxford: Clarendon Press.
 Dirac, P. A. M. 1936 *Proc. Roy. Soc. A*, **155**, 447.
 Fierz, M. & Pauli, W. 1939 *Proc. Roy. Soc. A*, **173**, 211.
 Harish-Chandra 1947 *Phys. Rev.* **71**, 793.
 Kemmer, N. 1939 *Proc. Roy. Soc. A*, **173**, 91.
 Kramers, H. A., Belinfante, F. J. & Lubanski, J. K. 1941 *Physica*, **8**, 597.
 Le Couteur, K. J. 1949 *Proc. Roy. Soc. A*, **196**, 251.
 Pauli, W. 1940 *Phys. Rev.* **58**, 716.
 Petiau, G. 1946 *J. Phys. théor. appl.* **6**, 181.
 Van der Waerden, B. L. 1932 *Gruppentheoretische methode in der quantenmechanik*. Berlin: Springer.
 Van Wyk, C. B. 1949 *Proc. Phys. Soc. A*, **62**, 697.
 Weyl, H. 1931 *Group theory and quantum mechanics*. London: Methuen.
 Wigner, E. 1939 *Ann. Math. Princeton, etc.*, **40**, 149.
 Wild, E. 1947 *Proc. Roy. Soc. A*, **191**, 253.

On the stochastic theory of continuous parametric systems and its application to electron cascades

BY H. J. BHABHA, F.R.S.

Tata Institute of Fundamental Research, Bombay

(Received 11 July 1949—Revised 31 March 1950)

A mathematical definition of an assembly of continuous parametric systems is given and its theory developed which makes it correspond precisely to most of the continuous parametric assemblies known in nature. Certain general theorems (formulae (19) and (22)) are deduced which hold for all such assemblies. It is shown that the usual method of treating a continuous parametric assembly by dividing up the domain of the parameter into a number of small segments, treating each segment as belonging to a discrete state of the system and then passing to the limit of making the segments infinitely small does not lead to a continuous parametric assembly of the type described above, but one of much wider generality which does not correspond to any type of physical system met with in nature. The general method and the theorems are of immediate application in calculating the fluctuations of the number of particles in chain reacting systems, and not only for systems in thermodynamic equilibrium.

The general theory is applied, as an illustration, to the stochastic treatment of an electron cascade to derive the differential equations which determine the functions from which the mean number of particles in any energy interval and the mean square deviation of this number can be calculated. It is shown in the appendix how the application of the usual method to this problem leads to the same results only if particular boundary conditions are imposed on the problem.

1. INTRODUCTION

When a physical system only possesses a number of discrete states, the stochastic treatment of an assembly of such systems presents no difficulties in principle. For brevity, it is convenient to call an assembly of such systems possessing only a number of discrete states an 'assembly of discrete state systems', or even a 'discrete state assembly'. One obtains all the statistical information that is possible about the assembly as soon as one can calculate the mean of the number of systems in any given state, the mean of the square and higher powers of this number, and the mean of all products such as the number in one state times the number in another state, etc.

In nature, however, there is a whole host of physical systems whose states pass continuously into one another, so that a state is labelled by giving a particular set of values to a set of parameters each of which can vary continuously in some domain. For example, a free material particle is of this type, and any of its states is labelled by giving a particular set of values to the three components of its momentum. We call an assembly of such systems a 'continuous parametric assembly' for brevity. The stochastic treatment of a continuous parametric assembly presents some interesting features. The usual method of treating such an assembly is to divide the domain of every parameter into a very large number of small segments and to assume that with regard to each segment there are only two possibilities for a system: either to lie in the segment, or to lie out of it. No difference is assumed to exist between the states of a system lying in different parts of the same segment. This

procedure amounts to substituting a discrete state system for the original continuous parametric system, and assuming that the former can be made to approximate to the latter to any desired degree of accuracy by making the segments sufficiently small. While this procedure may suffice in many practical problems, from the theoretical point of view it is incorrect to suppose that a continuous parametric assembly can be approached through a limiting process in this manner from an assembly of discrete state systems. The continuous parametric assemblies with which one is usually concerned possess certain properties which cannot be deduced in general by considering them as limits of a general assembly of discrete state systems, and these properties actually make their treatment simpler.

As an example, consider a system whose states are labelled by only one continuous parameter x with a domain extending from 0 to a say. We divide the domain of x into a large number, say ν , of small segments extending respectively from 0 to x_1 , x_1 to x_2 , x_2 to x_3 , and so on, and use the integers 1, 2, 3, ... to label these segments. These segments need not be equal, but for simplicity it is convenient to take them all equal to δ , where δ is a number which may be as small as we please. If n_i denotes the actual number of particles in the i th segment, then the state of the assembly is completely determined by giving a set of integral values to the complete set of ν numbers $n_1, n_2, \dots, n_\nu$. With every such state $(n_1, n_2, \dots, n_\nu)$ must be associated a real number $\pi(n_1, n_2, \dots, n_\nu)$ giving the probability of the assembly being in that state, and from a knowledge of all such probabilities one can calculate every statistic of the assembly. For example, one defines the mean of the product $(n_i)^r (n_j)^s \dots$ through the equation

$$\overline{(n_i)^r (n_j)^s \dots} = \sum (n_i)^r (n_j)^s \dots \pi(n_1, \dots, n_\nu), \quad (1)$$

where the sign Σ always stands for a sum over all integral values from 0 to ∞ of every n_i independently. If $N_{ij} = \sum_{k=i}^j n_k$ is the sum of the number of particles in all states from i to j , then its mean is given by

$$\overline{N_{ij}} = \sum_{k=i}^j \overline{n_k}, \quad (2a)$$

$$\text{while} \quad \overline{(N_{ij})^2} = \sum_{k=i}^j \overline{n_k^2} + \sum_{k=i}^j \sum_{l=i+k}^j \overline{n_k n_l}. \quad (2b)$$

We are not concerned here with an assembly in thermodynamic equilibrium. There is nothing in the mathematical theory or in definition (1) which requires that $\overline{n_i^2} = \overline{n_i}$, even in the limit $\delta \rightarrow 0$, and a probability distribution can easily be given for which it is not so. For example, π could be zero for all states of the assembly except those in which there are two particles in some one state i only, with none in any other state. For these states one could take

$$\pi(0, 0, \dots, n_i = 2, 0, 0, \dots) = \delta \quad \text{for all } i.$$

Then one would have $\overline{n_i} = 2\delta$, $\overline{n_i^2} = 4\delta$, $\overline{n_i^2} \neq \overline{n_i}$ for all i , though the mean of the total number of particles is finite, even in the limit $\delta \rightarrow 0$.

Assuming that the variation of $\overline{n_k}$, $\overline{n_k^2}$, $\overline{n_k n_l}$ from one state to the next tends to zero in the limit $\delta \rightarrow 0$, one could replace the summations in (2) by integrations; thus

$$\overline{N(x_i, x_j)} = \int_{x_i}^{x_j} \overline{n}(x) dx, \quad (3a)$$

$$\overline{N(x_i, x_j)^2} = \int_{x_i}^{x_j} \overline{n^2}(x) dx + \int_{x_i}^{x_j} \int_{x_i}^{x_j} \overline{n_2(x, x')} dx dx', \quad (3b)$$

where the new functions $\overline{n}(x)$, $\overline{n^2}(x)$, $\overline{n_2}(x, x')$ are defined by

$$\overline{n}(x_i) = \lim_{\delta \rightarrow 0} \overline{n_i}/\delta, \quad \overline{n^2}(x_i) = \lim_{\delta \rightarrow 0} \overline{n_i^2}/\delta, \quad \overline{n_2}(x_i, x_j) = \lim_{\delta \rightarrow 0} \overline{n_i n_j}/\delta^2, \quad (3c)$$

assuming these limits to exist.* Since $\overline{n_i^2} \neq \overline{n_i}^2$, even in the limit $\delta \rightarrow 0$, in general $\overline{n^2}(x) \neq \overline{n}(x)^2$. In any actual problem, however, it may be possible to deduce that $\overline{n_i^2} = \overline{n_i}^2$ even for an assembly which is not in statistical equilibrium. For example, it is shown in the appendix that this can be deduced to hold for the electrons in an electron cascade in the limit $\delta \rightarrow 0$ as a consequence of the actual boundary conditions of the problem, but it cannot be deduced in general. *The crux of the matter is that the continuous parametric assembly which is arrived at as the limit of a discrete state assembly in the above way has far more general properties than most of the actual assemblies with which one has to deal in nature, and, in consequence, certain general properties which the latter assemblies possess are absent in the more general assemblies described above. The purpose of this paper is to define mathematically a continuous parametric assembly which corresponds exactly to the continuous parametric assemblies with which one has to deal. It will be proved in the next section that for such a continuous parametric assembly, satisfying the mathematical assumptions given below, the statement*

$$\overline{N(x_i, x_j)^2} = \overline{N(x_i, x_j)}^2 + \int_{x_i}^{x_j} \int_{x_i}^{x_j} \overline{n_2}(x, x') dx dx'$$

is always true, and constitutes a general theorem. In a continuous parametric assembly of this type there is in fact no function corresponding to $\overline{n^2}(x)$. Thus, an attempt to deduce the behaviour of the continuous parametric assemblies usually found in nature by considering them as the limits of discrete state assemblies in the above manner leads one to miss certain simplifying features possessed by all such assemblies.

The circumstances pointed out in the last paragraph show that it is of some importance to formulate a mathematical definition and treatment of a continuous parametric assembly of a type corresponding to those met with in nature. The attempt to treat a continuous parametric assembly in a manner analogous to the usual one outlined above for treating a discrete state assembly would require that the set of numbers $(n_1, n_2, \dots, n_\nu)$ determining a state of the assembly be replaced now by a function $n(x)$, and in order to deal with every possible state of the assembly $n(x)$ must be allowed to take on any positive integral value or zero at any point x , and any other positive integral value or zero at any other point x' , with no relation

* An actual physical example where this is so is given in the appendix.

between the values $n(x)$ and $n(x')$ however close x and x' may be. $n(x)$ has therefore to be a highly discontinuous function of x , while the probability $\pi(n(x))$ has to be a functional of this discontinuous function. Apart from the mathematical complications of dealing with functions of this type, it is easy to see that the assembly so described would have more general properties than any physical assembly with which one has had to deal in nature so far. For the continuous parametric system has an infinite number of states, and when the total number of systems in the assembly is finite, the only states of the assembly which play a part are those in which $n(x) = 0$ everywhere except at a finite number of points. To describe a state of the assembly it would therefore be sufficient to specify the discrete values of the parameter x at which $n(x) \neq 0$ while giving the value of $n(x)$ at each of these points. With every such state must be coupled a certain probability. Actually, even the assembly so described has more general properties than any physical assembly I am aware of, and as far as physical problems are concerned one has only to deal with assemblies in which in any state of the assembly $n(x) = 0$ everywhere except at a finite number of points $x_1, x_2, \dots$ of the parameter x where $n(x) = 1$. This leads directly to the required treatment, and it only remains to formulate the idea in the most convenient mathematical form.

One can approach the problem from another angle by considering the natural methods of describing the states of discrete state and continuous parametric assemblies. If one were to make an observation on a discrete state assembly one could come to know the number of systems in each state, so that the state of the assembly could be appropriately and completely described by giving the set of numbers $n_1, n_2, \dots, n_\nu$. With every such state of the assembly one associates a probability $\pi(n_1, \dots, n_\nu)$, and the rest of the theory follows naturally. On the other hand, if one were to make an observation on a continuous parametric assembly, one might find that there were, say, k systems present, the systems being in the states with the labels $x_1, x_2, \dots, x_k$ respectively. Such, for example, would be the result of taking a Wilson chamber photograph in a magnetic field of an electron cascade—one might find k electrons with the momenta $p_1, \dots, p_k$ respectively. Thus the appropriate and complete way of labelling the states of a continuous parametric assembly would be to state that there are, in some given circumstances, k systems present, the systems being in the states with the labels $x_1, \dots, x_k$ respectively. All states of the assembly would then be covered by allowing each of the values $x_1, \dots, x_k$ to cover the entire domain of x and to let k take on the non-negative integral values $0, 1, 2, 3, \dots$ to allow for the number of systems in the assembly being $0, 1, 2, \dots$. This is the central idea of this paper, and together with a precise mathematical definition of the properties of the probability Π yields at once a precise method of treating a continuous parametric assembly.

The general theory is given in the next section. In §3 the general theory is applied as an illustration to the stochastic treatment of electron cascades, while in the appendix is given the usual treatment of the same problem by dividing the domain of the energy of the electrons and photons into a large number of small intervals. This example illustrates quite clearly the points which have been mentioned above. We have here an exact treatment of a quantum mechanically determined chain

process which is not in statistical equilibrium and in which not only the mean values but the fluctuations can be worked out on the basis of the exact quantum-mechanical formulae for radiation loss and pair creation given by Bethe and Heitler. The equations so deduced are solved by using the Mellin transform in a paper with A. Ramakrishnan which will be published separately. In §4 collision loss is considered and the terms deduced which have to be added to the equations of the preceding section to incorporate its effect.

2. STOCHASTIC THEORY OF A CONTINUOUS PARAMETRIC ASSEMBLY

Consider a physical system whose states are labelled by one continuous parameter x with its domain extending from 0 to ∞ . There would be no essential change in the theory if the domain extended from $-\infty$ to ∞ or from one finite real number to another. For brevity, we call a system in the state labelled by the value x 'a system of label x '. An eigenstate of an assembly of such systems is defined as the state in which the assembly finds itself after an observation is performed on it, and is completely determined by giving the labels of all the systems comprising it after the observation. We take into consideration the possibility that the number of systems in the assembly may change. Moreover, since the systems are assumed to be indistinguishable from each other, we make the convention of giving the labels in ascending order of magnitude. Thus a possible eigenstate of the assembly is one in which there are k systems present with the labels $x_1, \dots, x_k$ and we denote this eigenstate by

$$(x_1, x_2, \dots, x_k) \quad (0 \leq x_1 \leq x_2 \leq \dots \leq x_k). \quad (4)$$

Denote by X_k the k -dimensional space of all eigenstates defined by (4). An eigenstate is denoted by a point in the space. Denote by S_k any Lebesgue measurable set in X_k . For $k = 0$ the space X_0 consists of just one point. Let $X = UX_k$ be the space which is a union* of all the spaces X_k ($k = 0, 1, 2, \dots$) and let $S = US_k$ be the union of any sequence, finite or infinite, of sets $S_k, S_l, S_m, \dots$, each being Lebesgue measurable in $X_k, X_l, X_m, \dots$ respectively. We now make the following

DEFINITION 1. *With every such point set S in X is associated a real number $\Pi(S)$ having the following properties:*

$$(a) \quad 0 \leq \Pi(S) \leq 1. \quad (5a)$$

(b) Π is a completely additive function, so that for any sequence, finite or infinite, of point sets $S^1, S^2, \dots$ in X ,

$$\Pi(S^1 + S^2 + \dots) = \Pi(S^1) + \Pi(S^2) + \dots, \quad (5b)$$

if S^m and S^n have no point in common when $m \neq n$.

(c) $\Pi(X_k)$ exists for every k and

$$\Pi(X) = \sum_{k=0}^{\infty} \Pi(X_k) = 1. \quad (5c)$$

* Every point in X is a point in one X_k , and conversely any point in any X_k is a point in X .

(d) The series $\sum_{k=0}^{\infty} k^n \Pi(X_k)$ converges for every $n \geq 0$. (5d)

(e) $\Pi(S_k)$ is an absolutely continuous function of measurable sets $S_k \subset X_k$. That is, if $S_k \subset X_k$ is a point set of Lebesgue measure* $V(S_k)$, then

$$\Pi(S_k) \rightarrow 0 \quad \text{as} \quad V(S_k) \rightarrow 0. \quad (5e)$$

$\Pi(S)$ can then be interpreted as the probability of the assembly being in any one of the eigenstates corresponding to a point in S . Π may, of course, also be a function of one or more variables determining the state of the assembly, as, for example, the time, distance, temperature, etc.

It follows from the properties (a) to (e) of the above definition that a function $\pi(x_1, \dots, x_k)$ exists almost everywhere in X_k such that

$$\Pi(S_k) = \int \dots \int \pi(x_1, \dots, x_k) dx_1 \dots dx_k; \quad (6a)$$

(5c) can therefore also be written in the form

$$\Pi_0 + \sum_{k=1}^{\infty} \int_0^{\infty} dx_k \int_0^{x_k} dx_{k-1} \dots \int_0^{x_1} dx_1 \pi(x_1, \dots, x_k) = 1, \quad (6b)$$

where Π_0 stands for $\Pi(X_0)$, the value of Π for the eigenstate in which no system is present.

$\pi(x_1, \dots, x_k) dx_1 \dots dx_k$ can be interpreted as the probability of the assembly being in the state in which there are k systems present, one with its label lying between x_1 and $x_1 + dx_1$, another with its label between x_2 and $x_2 + dx_2$ and so on. It should be noted that the value of π itself need not be less than 1.

Let p be some property of the assembly whose value for the eigenstate $(x_1, \dots, x_k)$ is $p(x_1, \dots, x_k)$. We assume that p is a continuous and bounded function in the space X_k . Then the integral

$$\bar{p}(S_k) = \int dx_k \dots \int dx_1 \pi(x_1, \dots, x_k) p(x_1, \dots, x_k) \quad (7a)$$

exists and is called the mean value of p for all states of the assembly lying in S_k . Then $\bar{p}(X_k)$ likewise exists, as also the sum

$$\bar{p} = \sum_{k=0}^{\infty} \bar{p}(X_k), \quad (7b)$$

called the mean of p for the actual state of the assembly.

In order to avoid the awkwardness of handling limits of integration of the type which occur in (6b) we first extend the space X_k and then extend the definition of Π and any other function defined in X_k over the extended space. Let X'_k be the space consisting of all points $(x_1, \dots, x_k)$ without imposition of the condition (4). If P be the operation which changes the numbers $1, 2, \dots, k$ into $i_1, i_2, \dots, i_k$, where i_1 to i_k are the same integers 1 to k arranged in some other order, then by $P(x_1, \dots, x_k)$ we mean the point $(x_{i_1}, \dots, x_{i_k})$, and by $P(S_k)$ the set of all points $P(x_1, \dots, x_k)$, where

* k -dimensional volume.

$(x_1, \dots, x_k)$ is a point of the set S_k . We extend the definition of Π throughout the whole space X'_k through the definition

$$\Pi(P(S_k)) = \Pi(S_k). \quad (8a)$$

We likewise extend the definition of any point function $p(x_1, \dots, x_k)$ through the equation

$$p(x_{i_1}, \dots, x_{i_k}) = p(x_1, \dots, x_k). \quad (8b)$$

It then follows that

$$\bar{p}(X'_k) = k! \bar{p}(X_k), \quad (9)$$

since by (e) of definition 1, the contribution to the left- or the right-hand sides which comes from any common boundary of the space X_k with the space $P(X_k)$ is zero. (7) can therefore be replaced by

$$\bar{p} = \Pi_0 p_0 + \sum_{k=1}^{\infty} \frac{1}{k!} \int_0^{\infty} dx_k \int_0^{\infty} dx_{k-1} \dots \int_0^{\infty} dx_1 \pi(x_1, \dots, x_k) p(x_1, \dots, x_k). \quad (10)$$

We now introduce the tooth function T defined by

$$T(x | x_1 | x') = \begin{cases} 0 & \text{for } x_1 < x, \\ 1 & \text{for } x \leq x_1 \leq x', \\ 0 & \text{for } x_1 > x. \end{cases} \quad (11)$$

$$\text{Then } T \text{ satisfies } T(x | x_1 | x') T(x | x_1 | x') = T(x | x_1 | x'). \quad (12)$$

If $N(x, x')$ denotes the number of systems of label between x and $x + dx$, then in the eigenstate $(x_1, \dots, x_k)$ this number is clearly

$$N(x | x_1, \dots, x_k | x') = \sum_{i=1}^k T(x | x_i | x'). \quad (13)$$

It follows from (11) that

$$\begin{aligned} N(x | x_1, \dots, x_k | x')^2 &= \sum_{i=1}^k T(x | x_i | x') + \sum_{i=1}^k \sum_{j=1 \neq i}^k T(x | x_i | x') T(x | x_j | x') \\ &= N(x | x_1, \dots, x_k | x') + \sum_{i=1}^k \sum_{j=1 \neq i}^k T(x | x_i | x') T(x | x_j | x'). \end{aligned} \quad (14)$$

One can now form the mean values of $N(x, x')$ and $N(x, x')^2$ by substituting (13) and (14) respectively in the definition (10). This yields first, after some elementary manipulation,

$$\overline{N(x, x')} = \sum_{k=1}^{\infty} \frac{1}{(k-1)!} \int_x^{x'} dx'' \int_0^{\infty} dx_{k-1} \dots \int_0^{\infty} dx_1 \pi(x_1, \dots, x_{k-1}, x''). \quad (15)$$

It follows by a well-known theorem (Titchmarsh 1939, § 10.83) that we can change the order of the x'' integration and summation to obtain

$$\overline{N(x, x')} = \int_x^{x'} dx'' \cdot \bar{n}(x''), \quad (16)$$

$$\text{where } \bar{n}(x) = \sum_{k=1}^{\infty} \frac{1}{(k-1)!} \int_0^{\infty} dx_{k-1} \dots \int_0^{\infty} dx_1 \pi(x_1, \dots, x_{k-1}, x), \quad (17)$$

since the series (15) converges on account of the definition of Π and every term of the series (17) is ≥ 0 for every x .

$\bar{n}(x)dx$ can be interpreted as the mean of the number of systems whose labels lie between x and $x+dx$, but there is strictly no proper physical quantity $n(x)$ whose mean $\bar{n}(x)$ can be considered to be. One might, for example, think of defining $n(x)$ as a quantity whose value in the state $(x_1, \dots, x_k)$ is

$$n(x | x_1, \dots, x_k) = \sum_{i=1}^k \delta(x - x_i), \quad (18)$$

where $\delta(x)$ is the well-known delta function of Dirac. Substitution of (18) into (10) does in fact lead to (17), but this is not sufficient. If $n(x)$ were a proper physical magnitude, one should be able to form the higher moments like $\overline{(n(x))^2}$. The definition (18) then fails, for the integral of terms like $(\delta(x - x_i))^2$ is infinite. $\bar{n}$ has therefore to be considered as a single symbol denoting a new function of x defined by (17).

Substituting (14) in (10),

$$\overline{N(x, x')^2} = \overline{N(x, x')} + \sum_{k=2}^{\infty} \frac{1}{(k-2)!} \int_x^{x'} dx''' \int_x^{x'} dx'' \int_0^{\infty} dx_{k-2} \dots \int_0^{\infty} dx_1 \pi(x_1, \dots, x_{k-2}, x'', x''').$$

Again the orders of the summation and the x'', x''' integrations can be interchanged on account of definition 1 and we arrive at

THEOREM 1. *If $N(x, x')$ denote the number of systems with labels between x and x' in an assembly of systems whose states are labelled by one continuous parameter x , then the mean of its square is given by*

$$\overline{N(x, x')^2} = \overline{N(x, x')} + \int_x^{x'} dx'' \int_x^{x'} dx''' \bar{n}_2(x'', x'''), \quad (19)$$

where $\bar{n}_2(x'', x''')$ is a function of two variables x'' and x''' defined by

$$\bar{n}_2(x'', x''') = \sum_{k=0}^{\infty} \frac{1}{k!} \int_0^{\infty} dx_k \dots \int_0^{\infty} dx_1 \pi(x_1, \dots, x_k, x'', x'''). \quad (20)$$

$\bar{n}_2(x, x') dx dx'$ is the correlation and can be interpreted as the mean of the number of systems whose labels lie between x and $x+dx$ times the number whose labels lie between x' and $x'+dx'$ provided $x \neq x'$. (19) is a result of considerable importance in the general theory. It could not have been deduced by considering the continuum as the limit of a discrete state assembly except by making special restricting assumptions. $\bar{n}(x)$ is the derivative of $\overline{N(x, x')}$ and exists almost everywhere. Since the second term on the right of (19) is of order $(\delta x)^2$ for $x' - x = \delta x \rightarrow 0$ because $\bar{n}_2(x, x')$ is Lebesgue integrable, it follows that

$$\lim_{\delta \rightarrow 0} \frac{\overline{N(x, x + \delta x)^2}}{\delta x} = \lim_{\delta x \rightarrow 0} \frac{\overline{N(x, x + \delta x)}}{\delta x}, \quad (21)$$

wherever both the limits exist. This result is a direct consequence of property (e) in definition 1 of Π , and is not true for more general types of continuous parametric assemblies in which the property (e) is not postulated. It corresponds in the discrete state assembly to the special case in which the relation $\bar{n}_i^2 = \bar{n}_i$ holds.

Since the number of systems in the label interval between x and $x + dx$ need not be independent of the number in the label interval between x' and $x' + dx'$, its mean $\overline{n_2(x, x')} dx dx' \neq \overline{n(x)} dx \overline{n(x')} dx'$ for $x \neq x'$. However, whenever these two expressions are equal, as they are when independence prevails, then (19) simply becomes $\overline{N(x, x')^2} = \overline{N(x, x')} + (\overline{N(x, x')})^2$, a relation which holds in particular for the Poisson distribution.

The theorem given above can be generalized without much difficulty. This generalization is more conveniently written with the help of a special notation. Let $[1], [2], [3], \dots$ written as affixes to a symbol or an integration sign represent different independent intervals of the domain of x . By the overlap of k intervals $[1], [2], \dots, [k]$ we mean the interval consisting of all points common to every one of these intervals and denote it by $[1, 2, \dots, k]$, written as an affix. By using (10) and (13) one can then deduce the more general

THEOREM 2. *The mean value of the product $N^{[1]}N^{[2]} \dots N^{[k]}$ for any state of an assembly of systems labelled by one continuous parameter can be written in the form*

$$\begin{aligned} \overline{N^{[1]}N^{[2]} \dots N^{[k]}} = & \int^{[1, 2, \dots, k]} dx \overline{n(x)} + \Sigma \int^{[1, 2, \dots, r]} dx_1 \int^{[r+1, \dots, k]} dx_2 \overline{n_2(x_1, x_2)} \\ & + \Sigma \int^{[1, 2, \dots, r]} dx_1 \int^{[r+1, \dots, s]} dx_2 \int^{[s+1, \dots, k]} dx_3 \overline{n_3(x_1, x_2, x_3)} \\ & + \dots + \int^{[1]} dx_1 \int^{[2]} dx_2 \dots \int^{[k]} dx_k \overline{n_k(x_1, x_2, \dots, x_k)}, \end{aligned} \quad (22)$$

$$\text{where} \quad \overline{n_k(x_1, \dots, x_k)} = \sum_{j=0}^{\infty} \frac{1}{j!} \int_0^{\infty} dx'_j \dots \int_0^{\infty} dx'_1 \pi(x'_1, \dots, x'_j, x_1, \dots, x_k). \quad (23)$$

The summation sign before each term denotes a summation over like terms in which the integers 1 to k are distributed in all possible different ways between the brackets affixed to the integrals, there being no difference made between the order of the integers in a bracket nor between the order of the brackets. Thus, for $k = 4$, integrals of the type of the second term are to be summed with the groupings of the integers in brackets given by

$$[123][4], [234][1], [341][2], [412][3], [12][34], [13][24], [14][23].$$

Formula (19) is a particular case of (22) in which there are only two equal intervals $[1] = [2]$. As another example, we deduce from (22) by taking three equal intervals $[1] = [2] = [3]$,

$$\begin{aligned} \overline{N(x, x')^3} &= \overline{N(x, x')} + 3 \int_x^{x'} dx_1 \int_x^{x'} dx_2 \overline{n_2(x_1, x_2)} + \int_x^{x'} dx_1 \int_x^{x'} dx_2 \int_x^{x'} dx_3 \overline{n_3(x_1, x_2, x_3)} \\ &= 3\overline{N(x, x')^2} - 2\overline{N(x, x')} + \int_x^{x'} dx_1 \int_x^{x'} dx_2 \int_x^{x'} dx_3 \overline{n_3(x_1, x_2, x_3)}. \end{aligned}$$

The mean value of any physical magnitude for the assembly can always be expressed in terms of the π 's, as (10) indicates, but in most cases the mean values can be expressed more simply with the help of one of the functions defined by (23). For

example, if p is a physical quantity whose value in the state $(x_1, \dots, x_k)$ is $\sum_{i=1}^k p(x_i)$, where $p(x)$ is a given function of x , then the expression (10) reduces to

$$\bar{p} = \int_0^\infty p(x) \bar{n}(x) dx.$$

This would be true of the total kinetic energy of the assembly. Similarly, if q is a quantity whose value in the state $(x_1, \dots, x_k)$ is $\sum_{i=1}^k \sum_{j=1 \neq i}^k q(x_i, x_j)$, where $q(x, x')$ is a known function of two variables x and x' , then

$$\bar{q} = \int_0^\infty \int_0^\infty q(x, x') \bar{n}_2(x, x') dx dx'.$$

Such would be the case if there was, for example, a potential energy between any pair of systems depending on their labels.

Once one knows from theorem 2 that the mean and higher moments of the number of particles in any given interval are determined solely by the functions defined by (23), it is quite simple, in most cases, to write down the equations determining these without having recourse to the equation for the probability π itself.

The general theory given above can be extended, without any difficulty, to systems whose states are labelled by several continuous parameters. It can be extended likewise without difficulty to assemblies of several different types of physical systems. One can then calculate in a similar way, for example, the mean of the number of systems of one type in a given label interval times the number of systems of another type in another label interval.

Since the method of labelling the states of a continuous parametric assembly which underlies the theory is a most general one, the question arises as to how one can generalize the theory of this section to embrace more general types of continuous parametric assemblies such as those suggested by an analogy with a discrete state assembly. The answer is that one has to drop the property (e) of definition 1. $\Pi(S_k)$ may now be different from zero even though the set S_k is of measure zero, as, for example, when the set S_k is a hyperplane in the space X_k . Equation (6a) is now no longer true, but it can be proved (the proof follows from the Lebesgue decomposition of an additive set function (Saks 1937, p. 35 (14.10)) that

$$\Pi(S_k) = \int \dots \int \pi(x_1, \dots, x_k) dx_1 \dots dx_k + \Pi(S'_k) + \Pi(S''_k) + \dots, \quad (6c)$$

where the function $\pi(x_1, \dots, x_k)$ exists almost everywhere as before, and $S'_k, S''_k, \dots$ is a sequence of sets of measure zero, each set being contained in S_k and having no point in common with any of the others. If S'_k is an l -dimensional hyperplane in S_k ($l \leq k$), then $\Pi(S'_k)$ can again be expressed as the sum of an l -dimensional integral of a point function which exists almost everywhere in S'_k plus the values of Π on a finite number of hyperplanes of measure zero in S'_k . Although (14) remains true, theorem 1 ceases to hold because the second term on the right of (14) makes a contribution to $\bar{N}(x, x')^2$, a part of which is expressible in the form of a double integral like the second term on

the right of (19), while another part is expressible as a single integral with respect to a function of one variable x . The former comes from the first integral on the right of (6c), while the latter comes from the other terms on the right of (6c), namely, the integrals over the hyperplanes of lesser dimensionality. Thus, $\overline{N(x, x')^2}$ is equal to a single integral from x to x' of a function which is not $\bar{n}(x)$, plus a double integral from x to x' of a function of two variables. The new function of the single variable x is appropriately denoted by $\overline{n^2}(x)$. Although $\lim_{\delta x \rightarrow 0} \overline{N(x, x + \delta x)^2} / \delta x$ exists almost everywhere, it is equal to $\overline{n^2}(x)$ and not $\bar{n}(x)$ as in (21).

To take an example, in the space X_2 the subset S'_2 of measure zero over which Π does not vanish may be the line $x_1 = x_2$. If $\Pi(S'_2)$ is an absolutely continuous function of S'_2 , then a point function $\pi'(x_1)$ exists almost everywhere such that

$$\Pi(S'_2) = \int \pi'(x_1) dx_1.$$

Thus, the more general type of continuous parametric assembly just described would correspond to the limiting case as $\delta \rightarrow 0$ of the usual discrete state assembly. But it does not correspond to the physical assemblies that exist in nature, the correct description of which is provided by imposing the mathematical restriction (e) of definition 1. This immediately leads to the regularities embodied in theorems 1 and 2.

3. APPLICATION TO ELECTRON CASCADES

As an example, the method of the preceding section will be employed now to treat the fluctuations in electron cascades, a problem which is of considerable importance in itself. Its treatment by the usual method is given briefly in the appendix. We have to deal here with three different physical entities: positrons, electrons and photons. The usual approximation will be made of assuming that the problem is one-dimensional, and that the angular spread of the particles can be neglected. t will be used to denote the thickness of material traversed measured in units characteristic of the cascade theory (cf., for example, Bhabha & Chakrabarty 1943). With this assumption, the state of any particle or photon is simply labelled by one continuous parameter E whose domain extends from 0 to ∞ , namely, its energy. An eigenstate of the whole assembly is now labelled by three sets of numbers, thus:

$$(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l).$$

This eigenstate is one in which there are j positrons present of energies $\epsilon_1^+, \dots, \epsilon_j^+$, k electrons of energies $\epsilon_1^-, \dots, \epsilon_k^-$, and l photons of energies $\epsilon_1, \dots, \epsilon_l$. A semicolon is placed after each set of variables referring to one type of physical entity. The totality of all eigenstates of the assembly is provided by taking all values of the ϵ^+ 's, ϵ^- 's and ϵ 's satisfying conditions of the type (4), and by letting the integers j , k and l take on all non-negative integral values independently of each other. $\pi(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l; t)$ is a function of these $j + k + l$ variables, and we extend its definition in the manner (9) to cover all values of these variables by

making it a symmetric function of the set of j variables ϵ^+ , and similarly for the sets ϵ^- and ϵ . It is, of course, a function of t . Equation (10) is now replaced by

$$\bar{p} = \sum_{j=0}^{\infty} \sum_{k=0}^{\infty} \sum_{l=0}^{\infty} \frac{1}{j! k! l!} \int (d\epsilon^+)^j (d\epsilon^-)^k (d\epsilon)^l \pi(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l) \times p(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l), \quad (25)$$

where $(d\epsilon^+)^j$ stands for $d\epsilon_1^+ \dots d\epsilon_j^+$ with a similar meaning for the other sets, and an integral sign without the limits indicated denotes an integration with respect to all these variables extending independently from 0 to ∞ .

$N^+(E, E')$ stands for the number of positions of energies $\geq E$ and $\leq E'$. In the eigenstate $(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l)$ its value is

$$N^+(E | \epsilon_1^+, \dots; \epsilon_1^-, \dots; \epsilon_1, \dots | E') = \sum_{i=1}^j T(E | \epsilon_i^+ | E'), \quad (26)$$

corresponding to (13). We shall often have occasion to put $E' = E + dE$, where dE is a small increment which can be made as small as we please. It is therefore convenient to introduce the abbreviations $N(E,)$ for $N(E, E + dE)$ and $T(E | \epsilon_i^+ |)$ for $T(E | \epsilon_i^+ | E + dE)$. Similar quantities $N^-(E, E')$ and $M(E, E')$ must be introduced representing the number of electrons and photons respectively whose energies lie between E and E' .

As is well known, a cascade is built up by the operation of only two processes, radiation emission and pair-creation, the formulae for which have been calculated by Bethe and Heitler. $R(E, E') dE'$ denotes the probability per unit distance that an electron or positron of energy E radiates a photon of energy $E - E'$ and is left with an energy lying between E' and $E' + dE'$, and $R'(E, E') dE'$ the probability per unit distance that a photon of energy E creates a pair whose electron (or positron) has an energy lying between E' and $E' + dE'$. The formulae are quite symmetric between electrons and positrons so that $R'(E, E') = R'(E, E - E')$. The effect of collision loss will be considered in the next section. For the moment we neglect it. The set of equations determining the change of π with thickness t can then be written down at once. It is

$$\begin{aligned} & \frac{d}{dt} \pi(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l; t) \\ &= \sum_{r=1}^j \sum_{s=1}^l \pi(\epsilon_1^+, \dots, \epsilon_r^+ + \epsilon_s, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_s, \dots, \epsilon_l; t) R(\epsilon_r^+ + \epsilon_s, \epsilon_r^+) \\ &+ \sum_{r=1}^k \sum_{s=1}^l \pi(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_r^- + \epsilon_s, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_s, \dots, \epsilon_l; t) \\ &+ \sum_{r=1}^j \sum_{s=1}^k \pi(\epsilon_1^+, \dots, \epsilon_r^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_s^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l, \epsilon_r^+ + \epsilon_s^-; t) \lambda'(\epsilon_r^+ + \epsilon_s^-, \epsilon_r^+) \\ &- \pi(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l; t) \left\{ \sum_{r=1}^j \int_0^{\epsilon_r^+} R(\epsilon_r^+, \epsilon) d\epsilon \right. \\ &\left. + \sum_{r=1}^k \int_0^{\epsilon_r^-} R(\epsilon_r^-, \epsilon) d\epsilon + \sum_{r=1}^l \int_0^{\epsilon_r} R'(\epsilon_r, \epsilon) d\epsilon \right\}. \end{aligned} \quad (27)$$

A symbol with an underline indicates that that particular symbol which occurs in the π on the left of the equation is to be omitted from the π in which the underline occurs. The first term gives the increase in π due to the emission of radiation by positrons, the next term the increase due to radiation by electrons, the third term the increase due to pair creation, while the last term gives the decrease in π due to all of these three processes. This is the master equation of the cascade theory, if collision loss is neglected, from which all information about the cascade can be derived.

It is convenient at this stage to introduce functions of the type

$$\overline{n_p^+ n_q^- m_r} (E_1^+, \dots, E_p^+; E_1^-, \dots, E_q^-; E_1, \dots, E_r; t),$$

where the first pE 's refer to positron variables, the next q to electron variables and the last r to photon variables. This fact is indicated by separating the different sets of variables by semicolons and writing p, q and r as suffixes to n^+ , n^- and m respectively. The expression before the bracket is to be looked on as one composite symbol, and not as a product of factors. These functions correspond to the functions given by (23) and are defined by

$$\begin{aligned} \overline{n_p^+ n_q^- m_r} (E_1^+, \dots, E_p^+; E_1^-, \dots, E_q^-; E_1, \dots, E_r; t) \\ = \sum_{j=0}^{\infty} \sum_{k=0}^{\infty} \sum_{l=0}^{\infty} \frac{1}{j! k! l!} \int (d\epsilon^+)^j (d\epsilon^-)^k (d\epsilon)^l \pi(E_1^+, \dots, E_p^+, \epsilon_1^+, \dots, \epsilon_j^+; \\ E_1^-, \dots, E_q^-, \epsilon_1^-, \dots, \epsilon_k^-; E_1, \dots, E_r, \epsilon_1, \dots, \epsilon_l; t). \end{aligned} \quad (28)$$

We make the convention of not writing a suffix if its value is 1, and of omitting the symbol n^+ , n^- or m altogether if the corresponding suffix p, q or r is zero.

To get the rate of change of $N^+(E,)$ with t , we multiply (27) by

$$N^+(E | \epsilon_1^+, \dots; \epsilon_1^-, \dots; \epsilon_1, \dots |) = \sum_{i=1}^j T(E | \epsilon_i^+ |)$$

divided by $j! k! l!$ integrate over all the ϵ 's from 0 to ∞ and sum over all j, k and l independently from 0 to ∞ . We thus get on the left-hand side $d\{\overline{n^+}(E; t) dE\}/dt$, where $\overline{n^+}(E; t)$ is defined by (28). A summation from 0 to ∞ is implied whenever the limits of summation are not indicated. The first term on the right-hand side now gives

$$\begin{aligned} \sum_{j,k,l} \frac{1}{j! k! l!} \int (d\epsilon^+)^j (d\epsilon^-)^k (d\epsilon)^l \left[\sum_{r=1}^j \sum_{s=1}^l \pi(\epsilon_1^+, \dots, \epsilon_r^+ + \epsilon_s, \dots, \epsilon_j^+; \right. \\ \left. \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_s, \dots, \epsilon_l; t) R(\epsilon_r^+ + \epsilon_s, \epsilon_r^+) \right] \left\{ \sum_{i=1}^j T(E | \epsilon_i^+ |) \right\}. \end{aligned}$$

Introducing the new variable $\epsilon_r^{+'} = \epsilon_r^+ + \epsilon_s$, remembering that

$$\int_0^{\infty} d\epsilon_r^+ \int_0^{\infty} d\epsilon_s \dots = \int_0^{\infty} d\epsilon_r^{+'} \int_0^{\epsilon_r^{+'}} d\epsilon_r^+ \dots, \quad (29)$$

and using the symmetry of π in ϵ_1^+ to ϵ_j^+ and ϵ_1 to ϵ_l , this becomes

$$\sum_{j,k,l} \frac{1}{(j-1)!k!(l-1)!} \int (d\epsilon^+)^{j-1} (d\epsilon_j^{+'}) (d\epsilon^-)^k (d\epsilon)^{l-1} \pi(\epsilon_1^+, \dots, \epsilon_{j-1}^+, \epsilon_j^{+'}; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_{l-1}; t) \int_0^{\epsilon_j^{+'}} d\epsilon_j^+ R(\epsilon_j^{+'}, \epsilon_j^+) \left\{ \sum_{i=1}^j T(E | \epsilon_i^+ |) \right\} \quad (30)$$

$$= \sum_{j,k,l} dE \int_0^\infty d\epsilon_j^{+'} \left\{ \frac{1}{(j-2)!k!(l-1)!} \int (d\epsilon^+)^{j-2} (d\epsilon^-)^k (d\epsilon)^{l-1} \pi(\epsilon_1^+, \dots, \epsilon_{j-2}^+, E, \epsilon_j^{+'}; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_{l-1}; t) \int_0^{\epsilon_j^{+'}} d\epsilon_j^+ R(\epsilon_j^{+'}, \epsilon_j^+) \right\} \quad (31a)$$

$$+ \sum_{j,k,l} dE \int_E^\infty d\epsilon_j^{+'} \left\{ \frac{1}{(j-1)!k!(l-1)!} \int (d\epsilon^+)^{j-1} (d\epsilon^-)^k (d\epsilon)^{l-1} \pi(\epsilon_1^+, \dots, \epsilon_{j-1}^+, \epsilon_j^{+'}; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_{l-1}; t) R(\epsilon_j^{+'}, E) \right\} \quad (31b)$$

$$= dE \int_0^\infty d\epsilon^+ \bar{n}_2^+(E, \epsilon^+; t) \int_0^{\epsilon^+} d\epsilon^{+'} R(\epsilon^+, \epsilon^{+'}) \quad (32a)$$

$$+ dE \int_E^\infty d\epsilon^+ \bar{n}^+(\epsilon^+; t) R(\epsilon^+, E). \quad (32b)$$

The first term in the curly brackets in the last term of (27) gives

$$- \sum_{j,k,l} \frac{1}{(j-1)!k!l!} \int (d\epsilon^+)^j (d\epsilon^-)^k (d\epsilon)^l \pi(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l; t) \times \int_0^{\epsilon_j^+} d\epsilon^+ R(\epsilon_j^+, \epsilon^+) \left\{ \sum_{i=1}^j T(E | \epsilon_i^+ |) \right\} \quad (33)$$

$$= -dE \int_0^\infty d\epsilon^+ \bar{n}_2^+(E, \epsilon^+; t) \int_0^{\epsilon^+} d\epsilon^{+'} R(\epsilon^+, \epsilon^{+'}) \quad (34a)$$

$$- dE \bar{n}^+(E; t) \int_0^E d\epsilon^+ R(E, \epsilon^+). \quad (34b)$$

(32a) just cancels (34a), so that we are only left with (32b) and (34b) as far as the contribution of the terms so far considered is concerned.

Dealing with the other terms in (27) in the same manner, one arrives without difficulty at the well-known differential equation of the cascade theory for the variation of the mean number of positrons of energy between E and $E + dE$, namely,

$$\frac{d}{dt} \bar{n}^+(E; t) = \int_E^\infty \bar{n}^+(\epsilon; t) R(\epsilon, E) d\epsilon - \bar{n}^+(E; t) \int_0^\infty R(E, \epsilon) d\epsilon + \int_E^\infty \bar{m}(\epsilon; t) R'(\epsilon, E) d\epsilon. \quad (35)$$

Due to the symmetry in positrons and electrons one also gets an equation exactly of the form (35) with $\bar{n}^-(E; t)$ substituted in place of $\bar{n}^+(E; t)$.

Similarly, multiplying (27) by $M(E | \epsilon_1^+, \dots, \epsilon_l^-, \dots, \epsilon_1, \dots |)$ and integrating and summing to get the mean of $M(E)$, namely, $\bar{m}(E; t) dE$, one gets the physically obvious equation

$$\frac{d}{dt} \bar{m}(E; t) = \int_E^\infty \bar{n}^+(\epsilon; t) R(\epsilon, \epsilon - E) d\epsilon + \int_E^\infty \bar{n}^-(\epsilon; t) R(\epsilon, \epsilon - E) d\epsilon - \bar{m}(E; t) \int_0^E R'(E, \epsilon) d\epsilon. \quad (36)$$

Writing $\bar{n}(E; t) = \bar{n}^+(E; t) + \bar{n}^-(E; t)$ to denote the mean number of positrons plus electrons per unit energy range, and adding the equations for $\bar{n}^+$ and $\bar{n}^-$, one can get as usual just two differential equations for the two unknown functions $\bar{n}(E; t)$ and $\bar{m}(E; t)$. These are just the well-known equations of the cascade theory for the mean number of particles and photons in any infinitesimal energy interval. These equations could have been written down directly by considering the physical meaning of the functions $\bar{n}$ and $\bar{m}$, as has been done in the past, but it is instructive to derive them from the master equation (27), as it shows how terms of the type (32a) and (34a) which arise from different terms in the master equation eventually drop out.

To obtain the equation for $\bar{n}_2^+(\epsilon, \epsilon'; t)$ we multiply (27) by

$$\sum_{r=1}^j \sum_{s=1 \neq r}^j T(E | \epsilon_r^+ |) T(E | \epsilon_s^+ |), \quad (37)$$

integrate over all ϵ 's and then sum over all j, k, l from 0 to ∞ . This gives

$$\frac{d}{dt} \bar{n}_2^+(E, E'; t) dE dE'$$

on the left. The first term on the right of (27) can be written in the form (30) with (34) in place of the term in curly brackets. In consequence it gives

$$dE dE' \int_0^\infty d\epsilon' \bar{n}_3^+(E, E', \epsilon'; t) \int_0^{\epsilon'} d\epsilon R(\epsilon', \epsilon) \quad (38a)$$

$$+ dE dE' \int_{E'}^\infty d\epsilon \bar{n}_2^+(E, \epsilon; t) R(\epsilon, E') + dE dE' \int_E^\infty d\epsilon \bar{n}_2^+(E', \epsilon; t) R(\epsilon, E). \quad (38b)$$

The contribution of the first term in curly brackets in the last term of (27) is again obtained by putting (37) in place of the term in curly brackets in (33). It is

$$-dE dE' \int_0^\infty d\epsilon \bar{n}_3^+(E, E'; \epsilon; t) \int_0^\epsilon d\epsilon' R(\epsilon, \epsilon') \quad (39a)$$

$$- \bar{n}_2^+(E, E'; t) \left\{ \int_0^E R(E, \epsilon) d\epsilon + \int_0^{E'} R(E', \epsilon) d\epsilon \right\}. \quad (39b)$$

(39a) again cancels (38a). Dealing with the other terms in a like manner we arrive at the equation

$$\begin{aligned} \frac{d}{dt} \bar{n}_2^+(E, E'; t) = & \int_{E'}^\infty d\epsilon \bar{n}_2^+(E, \epsilon; t) R(\epsilon, E') + \int_E^\infty d\epsilon \bar{n}_2^+(E', \epsilon; t) R(\epsilon, E) \\ & + \int_{E'}^\infty d\epsilon \bar{n}^+ \bar{m}(E; \epsilon; t) R'(\epsilon, E') + \int_E^\infty d\epsilon \bar{n}^+ \bar{m}(E'; \epsilon; t) R'(\epsilon, E) \\ & - \bar{n}_2^+(E, E'; t) \left\{ \int_0^E R(E, \epsilon) d\epsilon + \int_0^{E'} R(E', \epsilon) d\epsilon \right\}. \end{aligned} \quad (40)$$

The physical meaning of all the terms on the right-hand side is self-evident.

The differential equation for $\overline{n^+m}(E; E'; t)$ which occurs in (40) is calculated in the same manner by multiplying the master equation (27) by

$$\sum_{r=1}^j \sum_{s=1}^l T(E | \epsilon_r^+ |) T(E' | \epsilon_s |), \quad (41)$$

and integrating and summing over all values of j, k and l . The contribution to this equation of all the terms in (27) can be calculated as before. The first term, however, yields a new feature in the equation, and it is therefore instructive to carry through the calculation here. The contribution of the first term on the right of (27) is given by (30) with (41) substituted in place of the expression in curly brackets in (30), and is

$$\begin{aligned} & \sum_{j,k,l} \frac{1}{(j-1)! k! (l-1)!} \int (d\epsilon^+)^{j-1} d\epsilon_j^{+'} (d\epsilon^-)^k (d\epsilon)^{l-1} \\ & \times \pi(\epsilon_1^+, \dots, \epsilon_{j-1}^+, \epsilon_j^{+'}; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_{l-1}; t) \\ & \times \int_0^{\epsilon_j^{+'}} d\epsilon_j^+ R(\epsilon_j^{+'}, \epsilon_j^+) \left\{ \sum_{r=1}^j \sum_{s=1}^l T(E | \epsilon_r^+ |) T(E' | \epsilon_s |) \right\}. \end{aligned}$$

The only special point to be remembered now is that $\epsilon_j^{+'} = \epsilon_j^+ + \epsilon_l$, so that the term $T(E' | \epsilon_s |)$ in curly brackets affects the result of the $\epsilon_j^{+'}$ integration. One gets, on summing over all j, k and l and dividing by $dE dE'$,

$$\int_0^\infty d\epsilon \overline{n_2^+m}(E, \epsilon; E'; t) \int_0^\epsilon d\epsilon' R(\epsilon, \epsilon') \quad (42a)$$

$$+ \int_E^\infty d\epsilon \overline{n^+m}(\epsilon; E'; t) R(\epsilon, E) \quad (42b)$$

$$+ \int_{E'}^\infty d\epsilon \overline{n_2^+}(E; \epsilon; t) R(\epsilon, \epsilon - E') \quad (42c)$$

$$+ \overline{n^+}(E + E'; t) R(E + E', E). \quad (42d)$$

(42c) and (42d) are directly the result of the circumstance just mentioned. (42d) is a correlation term of a type which does not occur in (40). It is the result of the circumstance that a positron of energy $E + E'$ can by the emission of radiation give rise simultaneously to a positron of energy E and a photon of energy E' , thus affecting the function $\overline{n^+m}(E; E'; t)$. No such term occurs in (40) because neither radiation loss nor pair creation leads to the simultaneous creation of two positrons. After some calculation one gets the equation

$$\begin{aligned} \frac{d}{dt} \overline{n^+m}(E; E'; t) &= \int_{E'}^\infty d\epsilon \overline{n_2^+}(E, \epsilon; t) R(\epsilon, \epsilon - E') + \int_{E'}^\infty d\epsilon \overline{n^+n^-}(E; \epsilon; t) R(\epsilon, \epsilon - E') \\ &+ \int_E^\infty d\epsilon \overline{n^+m}(\epsilon; E'; t) R(\epsilon, E) + \int_E^\infty d\epsilon \overline{m_2}(E', \epsilon; t) R'(\epsilon, E) \\ &- \overline{n^+m}(E, E'; t) \left\{ \int_0^E R(E, \epsilon) d\epsilon + \int_0^{E'} R'(E', \epsilon) d\epsilon \right\} + \overline{n^+}(E + E'; t) R(E + E', E). \end{aligned} \quad (43)$$

The physical meaning of all the terms in (43) is self-evident.

Owing to the symmetry of the theory in positrons and electrons one obtains the equations for $\overline{n_2^+}$ and $\overline{n_2^-}$ by replacing the index $+$ by $-$ in (40) and (43). The set of equations for the second-order correlation functions is completed by the equations for $\overline{n^+n^-}(E; E'; t)$ and $\overline{m_2}(E, E'; t)$, which can be calculated in the same way. They are

$$\begin{aligned} \frac{d}{dt} \overline{n^+n^-}(E; E'; t) = & \int_E^\infty d\epsilon \overline{n^+n^-}(\epsilon; E'; t) R(\epsilon, E) + \int_{E'}^\infty d\epsilon \overline{n^+n^-}(E; \epsilon; t) R(\epsilon, E') \\ & + \int_{E'}^\infty d\epsilon \overline{n^+m}(E; \epsilon; t) R'(\epsilon, E') + \int_E^\infty d\epsilon \overline{n^-m}(E'; \epsilon; t) R'(\epsilon, E) \\ & - \overline{n^+n^-}(E; E'; t) \left\{ \int_0^E R(E, \epsilon) d\epsilon + \int_0^{E'} R(E', \epsilon) d\epsilon \right\} + \overline{m}(E + E'; t) R'(E + E', E), \end{aligned} \quad (44)$$

$$\begin{aligned} \frac{d}{dt} \overline{m_2}(E, E'; t) = & \int_E^\infty d\epsilon \overline{n^+m}(\epsilon; E'; t) R(\epsilon, \epsilon - E) + \int_{E'}^\infty d\epsilon \overline{n^+m}(\epsilon; E; t) R(\epsilon, \epsilon - E') \\ & + \int_E^\infty d\epsilon \overline{n^-m}(\epsilon; E'; t) R(\epsilon, \epsilon - E) + \int_{E'}^\infty d\epsilon \overline{n^-m}(\epsilon; E; t) R(\epsilon, \epsilon - E') \\ & - \overline{m_2}(E, E'; t) \left\{ \int_0^E R'(E, \epsilon) d\epsilon + \int_0^{E'} R'(E', \epsilon) d\epsilon \right\}. \end{aligned} \quad (45)$$

The last term in (44) is again a correlation term which is due to the circumstance that a photon of energy $E + E'$ can create a pair whose positron has the energy E and the electron an energy E' , thus affecting the function $\overline{n^+n^-}(E; E'; t)$.

It should be noted in connexion with the equations (40) and (43) to (45) that while the functions $\overline{n_2^+}(E, E'; t)$, $\overline{n_2^-}(E, E'; t)$ and $\overline{m_2}(E, E'; t)$ are symmetric in E and E' , the other three functions $\overline{n^+n^-}$, $\overline{n^+m}$ and $\overline{n^-m}$ are not so.

The two equations of the type (35) together with (36) completely determine the three functions $\overline{n^+}$, $\overline{n^-}$ and $\overline{m}$ giving the mean number of positrons, electrons and photons respectively. The six equations of the type (40) and (43) to (45) then suffice to determine completely the six functions $\overline{n_2^+}$, $\overline{n^+n^-}$, $\overline{n_2^-}$, $\overline{n^+m}$, $\overline{n^-m}$ and $\overline{m_2}$. With the help of these the fluctuations in the number of positrons, electrons or photons in any energy interval can be calculated through formulae of the type (19), as also all second-order correlations such as the mean of the number of positrons in one energy interval times the number of electrons in another energy interval.

It has already been pointed out that $\overline{n_2^+}(E, E'; t) dE dE'$ represents the mean of the number of positrons in the energy interval from E to $E + dE$ times the number in the energy interval between E' and $E' + dE'$ provided the two energy intervals do not overlap, that is, provided $E \neq E'$ for the case $dE \rightarrow 0$, $dE' \rightarrow 0$. A similar interpretation holds for the functions $\overline{n_2^-}(E, E'; t)$ and $\overline{m_2}(E, E'; t)$. As a result of this interpretation one can write down without any calculation the integro-differential equations (40) and (43) to (45) which these functions must satisfy. The only possible doubt might at first sight be on the form of these equations at the point $E = E'$, for here the physical meaning of the functions, on the basis of which the equations can be written down, does not hold. However, it is clear from the continuity of these functions, based on

the continuity of π , that no additional terms can occur in the equations at the point $E = E'$. In case of doubt, however, one has always to refer back to equation (27) and deduce the relevant equations from it.

Introducing the functions

$$\overline{n_2}(E, E'; t) = \overline{n_2^+}(E, E'; t) + \overline{n_2^-}(E, E'; t) + \overline{n^+n^-}(E; E'; t) + \overline{n^+n^-}(E'; E; t), \quad (46)$$

$$\text{and} \quad \overline{nm}(E; E'; t) = \overline{n^+m}(E; E'; t) + \overline{n^-m}(E; E'; t), \quad (47)$$

it is easy to see that the six equations of the type (40) and (43) to (45) for the six functions $\overline{n_2^+}$, $\overline{n_2^-}$, $\overline{n^+n^-}$, $\overline{n^+m}$, $\overline{n^-m}$ and $\overline{m_2}$ can be combined into three integro-differential equations for the three independent functions $\overline{n_2}$, $\overline{nm}$ and $\overline{m_2}$.

A measure of the fluctuations in the mean number of particles in any interval is provided by calculating the mean of the square of the number of particles in the interval. Writing $N(E, E') = N^+(E, E') + N^-(E, E')$ for the number of particles, that is, positrons plus electrons, in the energy interval between E and E' ,

$$\overline{N(E, E')^2} = \overline{N^+(E, E')^2} + \overline{N^-(E, E')^2} + 2\overline{N^+(E, E')N^-(E, E')}. \quad (48)$$

Using the value of $N^+(E, E')$ given by (26) and the corresponding expression for $N^-(E, E')$, it is easy to see that

$$\overline{N^+(E, E')N^-(E, E')} = \int_E^{E'} d\epsilon \int_E^{E'} d\epsilon' \overline{n^+n^-}(\epsilon; \epsilon'; t), \quad (49)$$

where $\overline{n^+n^-}(\epsilon, \epsilon'; t)$ is given by (28). $\overline{N^+(E, E')^2}$ and $\overline{N^-(E, E')^2}$ are given by expressions corresponding to (19) with $\overline{n_2^+}$ and $\overline{n_2^-}$ again given by (28). We can, therefore, write

$$\overline{N(E, E')^2} = \overline{N(E, E')} + \int_E^{E'} d\epsilon' \int_E^{E'} d\epsilon \overline{n_2}(\epsilon, \epsilon'; t), \quad (50)$$

where $\overline{n_2}(\epsilon, \epsilon'; t)$ is given by (46).

It is important to realize that the last terms in (43) and (44) are essential for the correlation, for in their absence the equations (40) and (43) to (45) would be satisfied exactly by the substitution

$$\left. \begin{aligned} \overline{n_2^+}(E, E'; t) &= \overline{n^+}(E; t) \overline{n^+}(E'; t); & \overline{n_2^-}(E, E'; t) &= \overline{n^-}(E; t) \overline{n^-}(E'; t); \\ \overline{n^+n^-}(E; E'; t) &= \overline{n^+}(E; t) \overline{n^-}(E'; t); & \overline{m_2}(E, E'; t) &= \overline{m}(E; t) \overline{m}(E'; t); \\ \overline{n^+m}(E; E'; t) &= \overline{n^+}(E; t) \overline{m}(E'; t); & \overline{n^-m}(E; E'; t) &= \overline{n^-}(E; t) \overline{m}(E'; t). \end{aligned} \right\} \quad (51)$$

The number of particles or photons in one energy interval would then be independent of the number in any other energy interval. In particular, the last term in (50) would simply become $(\overline{N(E, E')})^2$, leading to the Poisson formula

$$\overline{N(E, E')^2} = \overline{N(E, E')} + (\overline{N(E, E')})^2. \quad (52)$$

This is, however, *not* so here, and the effect of the correlation terms is to give a fluctuation different from that to be expected on the Poisson formula.

4. EFFECT OF COLLISION LOSS

It only remains to find the terms which must be added to the equations of the preceding section to include the effect of collision loss on the behaviour of the cascade. As usual, we assume the collision loss to be a constant energy loss β per unit length.

The effect of collision loss is now twofold. First, it reduces the energy of every particle by an amount $\beta \delta t$ in traversing a distance δt . Secondly, a particle whose energy has already come to, or is, zero is then removed from the cascade altogether and travels no farther. These two effects have to be incorporated in the equations. Consider a state $(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l)$ of the assembly at the point $t + \delta t$. If no other effect were in operation, every particle in this state would have had an energy at the point t greater by an amount $\beta \delta t$, so that this state at $t + \delta t$ would be connected with the state $(\epsilon_1^+ + \beta \delta t, \dots, \epsilon_j^+ + \beta \delta t; \epsilon_1^- + \beta \delta t, \dots; \epsilon_1, \dots)$ at t . Secondly, any particle whose energy was less than $\beta \delta t$ at t would certainly be removed from the cascade before the point $t + \delta t$, so that the first particle would disappear from all states of the type $(x, \epsilon_1^+ + \beta \delta t, \dots, \epsilon_j^+ + \beta \delta t; \epsilon_1^- + \beta \delta t, \dots; \epsilon_1, \dots)$ with $x < \beta \delta t$ in going from t to $t + \delta t$. This state would then become the state $(\epsilon_1^+, \dots; \epsilon_1^-, \dots; \epsilon_1, \dots)$ at $t + \delta t$. The probability of the assembly being in all such states is

$$\int_0^{\beta \delta t} dx \pi(x, \epsilon_1^+ + \beta \delta t, \dots; \epsilon_1^- + \beta \delta t, \dots; \epsilon_1, \dots; t).$$

We therefore get, due to collision loss alone,

$$\begin{aligned} \pi(\epsilon_1^+, \dots, \epsilon_j^+; \epsilon_1^-, \dots, \epsilon_k^-; \epsilon_1, \dots, \epsilon_l; t + \delta t) &= \pi(\epsilon_1^+ + \beta \delta t, \dots; \epsilon_1^- + \beta \delta t, \dots; \epsilon_1, \dots; t) \\ &+ \int_0^{\beta \delta t} dx \pi(x, \epsilon_1^+ + \beta \delta t, \dots; \epsilon_1^- + \beta \delta t, \dots; \epsilon_1, \dots; t) \\ &+ \int_0^{\beta \delta t} dx \pi(\epsilon_1^+ + \beta \delta t, \dots; x, \epsilon_1^- + \beta \delta t, \dots; \epsilon_1, \dots; t). \end{aligned}$$

The last two terms are of order $\beta \delta t$. Terms of this type corresponding to states in which the number of particles is two or more greater than in the state on the left of this equation may be omitted, since their corresponding double or multiple integrals from 0 to $\beta \delta t$ would be of order $(\beta \delta t)^2$ or higher and would vanish in the limit $\delta \rightarrow 0$. Subtracting $\pi(\epsilon_1^+, \dots; \epsilon_1^-, \dots; \epsilon, \dots)$ from both sides of (73), dividing by δt and passing to the limit $\delta t \rightarrow 0$, we get on the left just the term on the left-hand side of (27), while on the right one gets

$$\left(\sum_{i=1}^j \frac{\partial}{\partial \epsilon_i^+} + \sum_{i=1}^k \frac{\partial}{\partial \epsilon_i^-} \right) \pi(\epsilon_1^+, \dots; \epsilon_1^-, \dots; \epsilon_1, \dots; t) \quad (53a)$$

$$+ \pi(0, \epsilon_1^+, \dots; \epsilon_1^-, \dots; \epsilon_1, \dots; t) \quad (53b)$$

$$+ \pi(\epsilon_1^+, \dots; 0, \epsilon_1^-, \dots; \epsilon_1, \dots; t). \quad (53c)$$

These are the terms which have to be added to the right-hand side of the master equation (27) to include the effect of collision loss.

The contribution of these terms to (35) is simply got by multiplying them by $\left\{ \sum_{i=1}^j T(E | \epsilon_i^+ |) \right\} / j! k! l!$ and integrating over the whole of space. Remembering the definition (28) we get at once from (53a) alone

$$\begin{aligned} dE \int_0^\infty d\epsilon \frac{\partial}{\partial \epsilon} \overline{n_2^+}(E, \epsilon; t) + \overline{n^+}(E + dE; t) - \overline{n^+}(E; t) + dE \int_0^\infty d\epsilon \frac{\partial}{\partial \epsilon} \overline{n^+ n^-}(E; \epsilon; t) \\ = dE \frac{\partial}{\partial E} \overline{n^+}(E; t) - dE \overline{n_2^+}(E, 0; t) - dE \overline{n^+ n^-}(E; 0; t), \end{aligned} \quad (54)$$

since $\overline{n_2^+}(E, \infty; t)$ and $\overline{n^+ n^-}(E; \infty; t)$ must be exactly zero for an assembly of finite energy. On the other hand, (53b) contributes just $\overline{n_2^+}(0, E; t)$, while (53c) contributes $\overline{n^+ n^-}(E; 0; t)$. These just cancel the second and third terms in (54). Thus the only term added to equation (35) is

$$\frac{\partial}{\partial E} \overline{n^+}(E; t), \quad (55)$$

while none is added to (36).

Only a simple calculation is required to derive that the terms to be added to equations (40), (43) and (44) are

$$\beta \frac{\partial}{\partial E} \overline{n_2^+}(E, E'; t) + \beta \frac{\partial}{\partial E'} \overline{n_2^+}(E, E'; t), \quad (56a)$$

$$\beta \frac{\partial}{\partial E} \overline{n^+ m}(E; E'; t), \quad (56b)$$

$$\text{and} \quad \beta \frac{\partial}{\partial E} \overline{n^+ n^-}(E; E'; t) + \beta \frac{\partial}{\partial E'} \overline{n^+ n^-}(E; E'; t) \quad (56c)$$

respectively, while none is to be added to (45).

APPENDIX

To give a stochastic treatment of electron cascades by the usual method one divides the domain of the energy E into a number of small segments extending from 0 to E_1 , E_1 to E_2 , E_2 to E_3 , etc., and labels these segments by the numbers 1, 2, 3, The segments need not be equal, but for simplicity it is convenient to take them equal to δ , where δ is a number which may be as small as we please. An eigenstate of the assembly is then described by giving the number of positrons, electrons and photons n_i^+ , n_i^- , and m_i respectively in the i th state for all values of i . With every such state of the assembly is associated a probability

$$\pi(n_1^+, n_2^+, \dots; n_1^-, n_2^-, \dots; m_1, m_2, \dots; t).$$

In the discrete case n_i^+ is a real physical magnitude and one can form not only its mean $\overline{n_i^+}$ but the higher moments $(\overline{n_i^+})^k$. These higher moments and all correlation terms are defined by a trivial generalization of definition (1). We now assume that the probability per unit distance of an electron or positron jumping from the state i to the state j with the emission of a photon is R_{ij} defined by

$$R_{ij} = R(E_i, E_j) \delta, \quad (57a)$$

while the corresponding probability for pair creation is

$$R'_{ij} = R'(E_i, E_j) \delta, \quad (57b)$$

where $R(E, E')$ and $R'(E, E')$ are the functions defined in §3. The equation for π is then

$$\begin{aligned} \frac{d}{dt} (n_1^+, \dots; n_1^-, \dots; m_1, \dots; t) \\ = - \left(\sum_i \sum_{r < i} n_i^+ R_{ir} + \sum_i \sum_{r < i} n_i^- R_{ir} + \sum_i \sum_{r < i} m_i R'_{ir} \right) \pi(n_1^+, \dots; n_1^-, \dots; m_1, \dots; t) \\ + \sum_i \sum_{r < i} (n_i^+ + 1) R_{ir} \pi(n_1^+, \dots, n_r^+ - 1, \dots, n_i^+ + 1, \dots; n_1^-, \dots; m_1, \dots, m_{i-r} - 1, \dots; t) \\ + \sum_i \sum_{r < i} (n_i^- + 1) R_{ir} \pi(n_1^+, \dots; n_1^-, \dots, n_r^- - 1, \dots, n_i^- + 1, \dots; m_1, \dots, m_{i-r} - 1, \dots; t) \\ + \sum_i \sum_{r < i} (m_i + 1) R'_{ir} \pi(n_1^+, \dots, n_r^+ - 1, \dots; n_1^-, \dots, n_{i-r}^- - 1, \dots; m_1, \dots, m_i + 1, \dots; t), \end{aligned} \quad (58)$$

where the dots stand for a given set of values of $n_1^+, n_2^+, \dots; n_1^-, n_2^-, \dots$ and $m_1, m_2, \dots$, and only those additional n 's and m 's are explicitly written which differ from those in the π on the left of the equation. Writing $n_i = n_i^+ + n_i^-$ for the total number of particles, i.e. electrons plus positrons in the i th state, one obtains from (58) without any difficulty the equation

$$\frac{d}{dt} \overline{n_i} = \sum_{j > i} \overline{n_j} R_{ji} - \overline{n_i} \left(\sum_{j < i} R_{ij} \right) + 2 \sum_{j > i} \overline{m_j} R'_{ji}, \quad (59)$$

which is just the obvious equation corresponding to the sum of (35) and the corresponding equation for electrons. In calculating the quadratic moments $\overline{n_i n_j}$ one has to distinguish between the cases in which $i \neq j$ and $i = j$, since the two can be and usually are of a different order. One obtains thus

$$\begin{aligned} \frac{d}{dt} \overline{n_i^2} = \sum_{j > i} \overline{n_j} R_{ji} + 2 \sum_{j > i} \overline{n_i n_j} R_{ji} - (2 \overline{n_i^2} - \overline{n_i}) \left(\sum_{j < i} R_{ij} \right) \\ + 2 \sum_{j < i} \overline{m_j} R'_{ji} + 4 \sum_{j > i} \overline{n_i m_j} R'_{ji} + 2 \overline{m_{2i}} R'_{2i, i}. \end{aligned} \quad (60)$$

By using (59) this simplifies to

$$\frac{d}{dt} (\overline{n_i^2} - \overline{n_i}) = -2(\overline{n_i^2} - \overline{n_i}) \left(\sum_{j < i} R_{ij} \right) + [2 \sum_{j > i} \overline{n_i n_j} R_{ji} + 4 \sum_{j > i} \overline{n_i m_j} R'_{ji} + 2 \overline{m_{2i}} R'_{2i, i}]. \quad (61)$$

On the other hand, for $i < j$, one obtains

$$\begin{aligned} \frac{d}{dt} \overline{n_i n_j} = \sum_{k > j} \overline{n_i n_k} R_{kj} + \sum_{k > i} \overline{n_k n_j} R_{ki} - \overline{n_i n_j} \left\{ \sum_{k < i} R_{ik} + \sum_{k < j} R_{jk} \right\} \\ + 2 \sum_{k > j} \overline{n_i m_k} R'_{kj} + 2 \sum_{k > i} \overline{n_j m_k} R'_{ki} + 2 \overline{m_{i+j}} R'_{i+j, i} + (\overline{n_j^2} - \overline{n_j}) R_{ji}. \end{aligned} \quad (62)$$

To proceed to the limit $\delta \rightarrow 0$ we note first that the orders of magnitude of $\overline{n_i}$, $\overline{n_i^2}$, $\overline{n_i n_j}$, etc., are determined not by the numerical factors n_i , n_i^2 , $n_i n_j$, etc., in the definition (1), but by the order of magnitude of the π 's involved. Secondly, the

number of states between two fixed energy values E and E' tends to infinity as $1/\delta$ when $\delta \rightarrow 0$. Hence, in order that the mean of the total number of particles in the energy interval between E and E' , and the mean of the square of this number shall be finite, it is necessary that $\bar{n}_i$ and $\bar{n}_i^2$ shall be of the order of magnitude δ , while $\bar{n}_i n_j$ ($i \neq j$) shall be of the order δ^2 as $\delta \rightarrow 0$, except in a finite number of intervals whose number remains finite as $\delta \rightarrow 0$. It follows that one can define the functions

$$\bar{n}(E_i) = \lim_{\delta \rightarrow 0} \bar{n}_i/\delta; \quad \bar{n}^2(E_i) = \lim_{\delta \rightarrow 0} \bar{n}_i^2/\delta; \quad \bar{n}_2(E_i; E_j) = \lim_{\delta \rightarrow 0} \bar{n}_i \bar{n}_j/\delta^2, \quad (63)$$

since the limits exist almost everywhere.* It follows that the expression in square brackets in (61) tends to zero as $\delta \rightarrow 0$, except for a possible contribution from the last term at a finite number of points at which $\bar{m}_{2i}$ is of order 1 instead of order δ . With the exception of such possible values of i , the number of which must always remain finite, (61) in the limit $\delta \rightarrow 0$ becomes almost everywhere

$$\frac{d}{dt} \{ \bar{n}^2(E) - \bar{n}(E) \} = -2 \{ \bar{n}^2(E) - \bar{n}(E) \} \int_0^E R(E, E') dE'. \quad (64)$$

For the actual boundary condition of the problem, corresponding to a shower started by a single electron or photon, or a spectrum of these, π is zero if any n_i or m_i is > 1 . In this case $\bar{n}^2(E) = \bar{n}(E)$ at $t = 0$, and equation (64) then shows that this relation holds for all t . But this condition is not mathematically necessary, and one could, for example, have at $t = 0$ a probability distribution defined by

$$\pi(0, 0, \dots, n_i^+ = 2, 0, 0, \dots) = \delta \quad \text{for all } i \text{ such that } E \leq E_i \leq E',$$

all other π 's zero. Then the limits (63) would exist, but we would have $\bar{n}^2(E) = 2\bar{n}(E)$, while the mean of the total number of particles would still be finite, namely, $\int_E^{E'} \bar{n}(E_i) dE_i$. As already pointed out in the text, this possibility is due to the fact that the continuous parametric assembly which results from a discrete state assembly by a limiting process in the above manner is of a more general nature than the one defined in §2 which corresponds to the assemblies actually found in nature. When $\bar{n}^2(E) = \bar{n}(E)$, as it is with the actual boundary condition, then the last term of (62) vanishes, and in this case in the limit $\delta \rightarrow 0$ (62) just becomes the equation for $\bar{n}_2(E, E')$ corresponding to (40) of the text.

I wish to thank Professor M. H. Stone, Professor D. D. Kosambi and Dr K. Chandrasekharan for useful discussions on the mathematics in the paper.

REFERENCES

- Bhabha, H. J. & Chakrabarty, S. K. 1943 *Proc. Roy. Soc. A*, **181**, 267–303.
 Saks, S. 1937 *Theory of the integral*. 2nd rev. ed. New York: G. E. Stechart and Co.
 Titchmarsh, E. C. 1939 *The theory of functions*, 2nd ed. Oxford University Press.

* That is, everywhere except over a point set of measure zero.

The molecular orbital theory of chemical valency.

V. The structure of water and similar molecules

BY J. A. POPLE, *Department of Theoretical Chemistry, University of Cambridge*

(Communicated by Sir John Lennard-Jones—Received 16 December 1949)

The theory of molecular and equivalent orbitals developed in previous papers of this series is used to discuss the spatial distribution of lone-pair electrons in molecules such as H_2O and NH_3 and the part they play in determining the equilibrium configuration. Previous treatments of H_2O have assumed that the lone pairs are essentially unaltered by molecular formation. It is shown here, on the other hand, that they will be displaced so as to be mainly concentrated on the side of the O-nucleus remote from the hydrogen atoms. An important consequence of this is that the lone-pair electrons will make a contribution to the total dipole moment. Comparison of the experimentally observed moment with an approximate quantitative treatment suggests that, as a result of this, transfer of electrons from the hydrogen atoms to the oxygen does not occur to the extent that has previously been believed.

The variation of the spatial distribution of the orbitals of H_2O with changes of nuclear configuration is examined and it is shown that, in the equilibrium position, the electronic structure can be described approximately by two sets of two equivalent orbitals pointing in nearly tetrahedral directions. The dependence of total energy on bond angle is discussed and it is shown that electrostatic repulsions between the equivalent orbitals are major factors in determining the equilibrium configuration. Similar considerations apply to NH_3 .

1. INTRODUCTION

Theoretical treatments of the structure of molecules with lone-pair electrons are usually based on the electronic configuration of the ground state of the corresponding atom. The most commonly discussed example is the water molecule (Van Vleck & Sherman 1935; Coulson 1941), the starting-point for which is the free oxygen atom with the configuration $(1s)^2(2s)^2(2p)^4$. Thus Van Vleck & Sherman suppose the orbitals $1s$, $2s$ and $2p_z$ to be doubly occupied, while $2p_x$ and $2p_y$ have only one electron each and are available for bonding. By examining the exchange energy between these latter orbitals and the hydrogen atomic functions, it is deduced that, subject to certain approximations, the bond angle will be 90° . The fact that the observed angle is larger is attributed partly to hydrogen repulsions and partly to hybridization with s -functions.

This treatment presupposes that the electrons represented by $(1s)^2(2s)^2(2p_z)^2$ do not take part directly in the bond formation and are substantially unaltered in the molecule. The analysis is, in fact, only valid if we impose the rigid restriction that the lone pairs must remain in these orbitals, a limitation which is likely to be very artificial in view of the powerful distorting forces due to the neighbouring hydrogen atoms. This restriction will mask the real causes preventing the molecule opening out into linear form.

In this paper the structure of molecules such as ammonia and water is examined by a different method, which attempts to avoid the undue use of the electronic structure of free atoms and which allows for the effect of molecular formation on *all* the electrons.

In the first three papers of this series (Lennard-Jones 1949 *a, b*; Hall & Lennard-Jones 1950) a general method for the description of the electrons of a molecule by a determinant of one-electron functions has been developed. In IV (Lennard-Jones & Pople 1950) the theory of inner shell and lone-pair electrons is discussed and it is shown how a transformation can be applied to the molecular orbitals to give sets of equivalent bond orbitals and sets of lone-pair orbitals. In H_2O , for example, there are two such sets, each with two members. In general there is an element of arbitrariness in the definition of these sets, owing to the possibility of forming them from a linear combination of molecular orbitals of the same symmetry. Suppose, for example, that ψ_1 and ψ_2 are orthogonal A_1 molecular orbitals which can be used in forming equivalent sets. Then two different sets, which are formally just as correct, can be obtained by applying the same transformation to the functions

$$\{(\cos \lambda) \psi_1 + (\sin \lambda) \psi_2\} \quad \text{and} \quad \{-(\sin \lambda) \psi_1 + (\cos \lambda) \psi_2\},$$

where λ is an arbitrary parameter. The existence of the parameter λ corresponds to a certain artificiality in our division of orbitals into bonds and lone pairs, but as the total wave function is invariant, we may fix it according to a suitable, convenient criterion of localization. As in the example of an atom in a uniform electric field discussed in IV, this leads to two sets of orbitals on opposite sides of the nucleus.

In this paper, the relative orientation of the bonding and lone-pair equivalent orbitals in the water molecule is examined in more detail, with a view to assessing the importance of lone-pair electrons in determining the observable properties of the molecule. One important consequence of the fact that the lone-pair equivalent orbitals are directed away from the hydrogen side of the molecule, is that they will make a considerable contribution to the total dipole moment. In dealing with properties connected with electronegativity it is not sufficient to discuss the transfer of charge to the more electronegative atom. The effect of the reorientation of the lone pairs must also be taken into account.

As pointed out in § 7 of IV, lone-pair electrons may be important in determining the equilibrium configuration. Here we shall discuss the importance of certain terms in the total energy expression involving lone pairs in relation to bond angles in molecules such as ammonia and water. It turns out that the interaction between lone pairs and between lone pairs and bond electrons are major factors preventing such molecules taking up their most symmetrical form.

2. THE DISTRIBUTION OF MOLECULAR AND EQUIVALENT ORBITALS IN H_2O AND NH_3

As it is not possible at present to calculate exact forms for molecular or equivalent orbitals, it is usual to use approximate analytical expressions. The most common of these (Lennard-Jones 1929; Mulliken 1935 *a*) is to use a linear combination of atomic orbitals. We shall use only the occupied atomic functions of the free atom as a simple approximation, although more accurate results could be obtained by allowing the orbitals to hybridize with higher states. Approximate molecular orbitals may be formed from functions based on all the atoms in the molecule, but for equivalent

orbitals, owing to their increased localization, functions centred on nuclei remote from the region of the orbital will not make a large contribution.

In the case of H_2O , the linear combination of atomic orbitals (approximate forms for the five molecular orbitals discussed in IV) are

$$\left. \begin{aligned} \omega_1(A_1) &= \phi(\text{O}; 1s), \\ \omega_2(B_1) &= b_2 \phi(\text{O}; 2p_y) + c_2 \{ \phi(\text{H}_{(1)}; 1s) - \phi(\text{H}_{(2)}; 1s) \}, \\ \omega_3(A_1) &= a_3 \phi(\text{O}; 2s) + b_3 \phi(\text{O}; 2p_x) + c_3 \{ \phi(\text{H}_{(1)}; 1s) + \phi(\text{H}_{(2)}; 1s) \}, \\ \omega_4(A_1) &= a_4 \phi(\text{O}; 2s) + b_4 \phi(\text{O}; 2p_x) + c_4 \{ \phi(\text{H}_{(1)}; 1s) + \phi(\text{H}_{(2)}; 1s) \}, \\ \omega_5(B_2) &= \phi(\text{O}; 2p_z), \end{aligned} \right\} \quad (2.01)$$

where the symbols in brackets on the left-hand sides describe their symmetry properties. $\phi(\text{O}; 1s)$, $\phi(\text{O}; 2s)$ and $\phi(\text{O}; 2p)$ are the oxygen $1s$, $2s$ and $2p$ atomic functions and $\phi(\text{H}_{(i)}; 1s)$ is the atomic hydrogen function associated with the proton i . The x -direction is chosen to be along the molecular axis and the z -direction perpendicular to the nuclear plane. We have omitted any contribution from the oxygen $1s$ functions to ω_3 and ω_4 as these are tightly bound to the nucleus and little disturbed by the external protons.

To simulate the properties of the true molecular orbitals in the sense of part III, the approximate forms for ω_3 and ω_4 (2.01) must be orthogonal and must satisfy the additional condition $E_{34} = 0$. This fixes their spatial distribution. It is probable that one, ω_3 say, will be mainly on the hydrogen side of the O, while the condition of orthogonality forces the other to the side of the O remote from the H-atoms. This means that the coefficient c_4 will be small.

However, unless we are concerned with ionization or similar processes, it is not necessary to deal with molecular orbitals. Instead, we may use two linear combinations of them, s_3 and s_4 , as discussed in the introduction. This means that we may relax the condition $E_{34} = 0$ so that s_3 and s_4 are symmetry orbitals according to the definition of part III. In the linear combination of atomic orbitals approximation, we shall fix the degree of freedom thereby introduced so that s_4 and hence the equivalent lone pairs are built out of functions centred only on the oxygen nucleus. This is merely a convenient way of fixing the meaning of the term lone pairs in this approximation. The symmetry orbitals s_3 and s_4 will have the form

$$\left. \begin{aligned} s_3 &= a'_3 \phi(\text{O}; 2s) + b'_3 \phi(\text{O}; 2p_x) + c'_3 \{ \phi(\text{H}_{(1)}; 1s) + \phi(\text{H}_{(2)}; 1s) \}, \\ s_4 &= a'_4 \phi(\text{O}; 2s) + b'_4 \phi(\text{O}; 2p_x). \end{aligned} \right\} \quad (2.02)$$

The bonding orbital s_3 is formed from the combined hydrogen function

$$\{ \phi(\text{H}_{(1)}; 1s) + \phi(\text{H}_{(2)}; 1s) \},$$

and an atomic directed hybrid orbital concentrated on the hydrogen side of the O-nucleus. Thus the coefficients a'_3 , b'_3 and c'_3 will be of the same sign. If we now define the (positive) overlap integrals

$$\left. \begin{aligned} S_1 &= \int \phi(\text{O}; 2s) \{ \phi(\text{H}_{(1)}; 1s) + \phi(\text{H}_{(2)}; 1s) \} d\tau, \\ S_2 &= \int \phi(\text{O}; 2p_x) \{ \phi(\text{H}_{(1)}; 1s) + \phi(\text{H}_{(2)}; 1s) \} d\tau, \end{aligned} \right\} \quad (2.03)$$

the condition of orthogonality between s_3 and s_4 leads to

$$\frac{a'_4}{b'_4} = -\frac{b'_3 + c'_3 S_2}{a'_3 + c'_3 S_1}. \quad (2.04)$$

This shows that a'_4 and b'_4 are of opposite sign and we infer that the orbital s_4 is concentrated on the remote side of the oxygen nucleus.

The equivalent lone pairs are formed by linear combinations of s_4 with the B_2 orbital ω_5 . As ω_5 is simply $\phi(\text{O}; 2p_z)$ in the first approximation, these lone pairs will be orientated away from the hydrogen side. It may be noted that a'_4/b'_4 is negative even if a'_3 or b'_3 vanishes. This means that the lone pair symmetry orbital s_4 is concentrated on the remote side even if there is no hybridization in the oxygen part of s_3 .

Equation (2.04), combined with the normalization condition, is sufficient to determine the coefficients in s_4 if those in s_3 are known. That is, the forms of the equivalent lone pairs are fixed by orthogonality conditions alone. This is due to the very limited flexibility allowed to the molecular orbitals by the restriction to the forms (2.01). The orbital ω_5 , for example, is put equal to $\phi(\text{O}; 2p_z)$ and is allowed no flexibility at all.

In order to decide whether the conclusion that the lone pairs are concentrated on the remote side of the oxygen nucleus is true of the exact functions, or whether it is merely a consequence of the crudeness of the forms (2.01), we have to examine the effect of higher approximations. The distortion is similar in kind to that of an atomic closed shell in a uniform electric field, as discussed briefly in part IV. In this case the orbitals will tend to be moved along with the field *as far as this is permitted by the orthogonality conditions*.

If higher approximations are used by allowing the orbitals to hybridize with other functions with more angular nodes, it will be possible to allow for the tendency of them all to move into regions of lower energy. For example, if the molecular orbital ω_5 in (2.01) is expressed in the more general form

$$\omega_5 = k\phi(\text{O}; 2p_z) + l\phi(\text{O}; 3d_{zz}), \quad (2.05)$$

variation of k and l to make the total energy a minimum will give a good idea of how much the lone-pair electron distribution drifts towards the protons. Similar considerations apply to the other orbitals.

Now for elements in the first row of the periodic table the $3d$ functions have considerably higher energies than $2s$ and $2p$ functions, so that coefficients such as l will be small. Thus we conclude that for molecules formed from these elements, the simple functions (2.01) leading to equation (2.04) give an essentially correct picture of the spatial distribution of the lone-pair orbitals, although the extent to which they are oriented away from the hydrogen side of the molecule is probably exaggerated. For similar elements in the next and higher rows, however, hybridization of the lone pairs with the corresponding functions will occur much more easily, as the energy difference between $3p$ and $3d$ functions is smaller than that between $2p$ and $3d$.

Returning to the linear combination of atomic orbitals approximations for H_2O , we are now in a position to discuss the formation of the equivalent bond orbitals from s_3 of equation (2.02) and the B_1 molecular orbital ω_2 .

In general the result will be

$$\left. \begin{aligned} \chi(b_1) &= d\phi(O; b_1^*) + \frac{1}{\sqrt{2}}(c'_3 + c_2)\phi(H_{(1)}; 1s) + \frac{1}{\sqrt{2}}(c'_3 - c_2)\phi(H_{(2)}; 1s), \\ \chi(b_2) &= d\phi(O; b_2^*) + \frac{1}{\sqrt{2}}(c_2 - c'_3)\phi(H_{(1)}; 1s) + \frac{1}{\sqrt{2}}(c'_3 + c_2)\phi(H_{(2)}; 1s), \end{aligned} \right\} \quad (2.06)$$

where $\phi(O; b_1^*)$ and $\phi(O; b_2^*)$ are certain oxygen atomic directed orbitals, not necessarily pointing towards the nuclei $H_{(1)}$ and $H_{(2)}$. We see that the equivalent bond orbitals given by (2.06) are not formed from atomic orbitals on two centres only. The terms with coefficients $(c'_3 - c_2)$ represent a certain delocalization. Coulson (1937, 1942*a*) has, in fact, found similar terms in the case of CH_4 . More recently he has discussed the importance of such terms in explaining the cross-terms in the vibrational potential function (Coulson 1948, 1949). However, these terms are smaller than the others, and as we are chiefly concerned with lone pairs, we will put $c_2 = c'_3$ in this approximate theory. We shall make the further approximation of supposing the directions b_1^* and b_2^* to coincide with the internuclear directions b_1 and b_2 .

We can now write the approximate forms for the two sets of equivalent orbitals in H_2O as

$$\left. \begin{aligned} \chi(b_1) &= \lambda\phi(O; b_1) + \mu\phi(H_{(1)}; 1s), \\ \chi(b_2) &= \lambda\phi(O; b_2) + \mu\phi(H_{(2)}; 1s), \end{aligned} \right\} \quad (2.07)$$

$$\text{and} \quad \left. \begin{aligned} \chi(l_1) &= \phi(O; l_1), \\ \chi(l_2) &= \phi(O; l_2), \end{aligned} \right\} \quad (2.08)$$

where the oxygen atomic orbitals are formed by hybridization of $2s$ and $2p$ functions:

$$\left. \begin{aligned} \phi(O; b_1) &= \cos \epsilon_b \phi(O; 2s) + \sin \epsilon_b \phi(O; 2p(b_1)), \\ \phi(O; b_2) &= \cos \epsilon_b \phi(O; 2s) + \sin \epsilon_b \phi(O; 2p(b_2)), \\ \phi(O; l_1) &= \cos \epsilon_l \phi(O; 2s) + \sin \epsilon_l \phi(O; 2p(l_1)), \\ \phi(O; l_2) &= \cos \epsilon_l \phi(O; 2s) + \sin \epsilon_l \phi(O; 2p(l_2)). \end{aligned} \right\} \quad (2.09)$$

Relations exist between the hybridization parameters ϵ_b , ϵ_l and the geometrical angles of the system. Let α be the bond angle, β the angle between the lone-pair directions l_1 and l_2 , and γ the angle between a bond and a lone pair. These are related geometrically by

$$\cos \gamma = -\cos \frac{1}{2}\alpha \cos \frac{1}{2}\beta. \quad (2.10)$$

The condition of orthogonality between $\chi(b_1)$ and $\chi(b_2)$ leads to

$$\lambda^2(\cos^2 \epsilon_b + \sin^2 \epsilon_b \cos \alpha) + 2\lambda\mu(S_s \cos \epsilon_b + S_p \sin \epsilon_b \cos \alpha) + \mu^2 S_H = 0, \quad (2.11)$$

where

$$\left. \begin{aligned} S_H &= \int \phi(H_{(1)}; 1s) \phi(H_{(2)}; 1s) d\tau, \\ S_s &= \int \phi(O; 2s) \phi(H; 1s) d\tau, \\ S_p &= \int \phi(O; 2p(b_1)) \phi(H_{(1)}; 1s) d\tau. \end{aligned} \right\} \quad (2.12)$$

The full orthogonality condition (2.11) is difficult to handle, so we shall make $\phi(O; b_1)$ and $\phi(O; b_2)$ orthogonal instead. This corresponds to taking $\mu = 0$ in (2.11) and leads to

$$\cot^2 \epsilon_b = -\cos \alpha. \quad (2.13)$$

For finite λ/μ it is found that the overlap integral $\int \chi(b_1) \chi(b_2) d\tau$ is reasonably small if (2.13) is satisfied. If λ/μ is 1.15, the value proposed in the next section, the overlap integral is 0.20.

Orthogonality between lone pairs leads to

$$\cot^2 \epsilon_l = -\cos \beta. \quad (2.14)$$

To obtain a relation between α and β we have to examine the condition of orthogonality between $\chi(b)$ and $\chi(l)$. This leads to

$$\frac{\lambda}{\mu} = -\frac{\cos \epsilon_l S_s + \sin \epsilon_l \cos \gamma S_p}{\cos \epsilon_b \cos \epsilon_l + \sin \epsilon_b \sin \epsilon_l \cos \gamma}. \quad (2.15)$$

This equation gives ϵ_l as a function of ϵ_b and λ/μ so that the lone pairs are determined when the bond orbitals are given.

We may note that, if the ratio λ/μ is given, the absolute values of λ and μ are determined by the normalization conditions for $\chi(b)$. Thus we have

$$1/\mu^2 = 1 + (\lambda/\mu)^2 + 2(\lambda/\mu)(S_s \cos \epsilon_b + S_p \sin \epsilon_b). \quad (2.16)$$

Equations (2.10) and (2.13) to (2.16) are now sufficient to determine these approximate forms for all the orbitals for a given bond angle α and a given value for the ratio λ/μ . We can get a rough idea of the way the orbitals vary as α increases from 90 to 180° by examining the particular case $\lambda/\mu = \infty$, which corresponds to a bond orbital centred only on the oxygen nucleus. In this case γ can be eliminated from (2.10) and (2.15) to give an equation for β when α is known, viz.

$$(1 + \sec \alpha)(1 + \sec \beta) = 4. \quad (2.17)$$

When $\alpha = 90^\circ$, the $\chi(b)$ are pure p -functions and the $\chi(l)$ diagonal hybrids, which point in opposite directions, each being perpendicular to the nuclear plane. As α increases β is reduced, so that the orbitals $\chi(l)$ are forced together. At $\alpha = \cos^{-1}(-\frac{1}{3})$ we have the case of four equivalent tetrahedral orbitals, and finally at $\alpha = 180^\circ$, the $\chi(l)$ are perpendicular p -functions. A more detailed study later in the paper shows that the orbitals behave similarly for finite λ/μ , the main difference being that the lone-pair orbitals are oriented further away from the hydrogen nuclei. The effect of these orbital changes on the total energy will be examined in § 4.

The theory for NH_3 is very similar to that for H_2O , except that we have to deal with three equivalent bond orbitals and only one lone pair. We shall not go into it in detail, but only indicate the orbital variation deduced by similar methods. If the bond angle is 90°, the lone-pair orbital will consist mainly of the nitrogen 2s function, but with a certain amount of 2p (if λ/μ is finite), causing it to be concentrated mainly on the remote side of the nucleus. As α is increased, an intermediate state will be reached, corresponding to an approximately tetrahedral configuration. Finally, if it reaches 120°, the four nuclei will lie in a plane and the lone-pair orbital will be a p -function distributed equally on both sides of the nucleus.

3. THE POLARITY OF THE BOND ORBITAL

The only factor left unspecified in the approximate wave functions which have been set up for H_2O is the ratio λ/μ . This should be determined by minimizing the total energy. Such calculations, however, are found to be very sensitive to the approximations that it is necessary to make, particularly for the kinetic energy terms, and it is not possible to get any reliable result. Instead λ/μ will be estimated by an empirical method. It should be mentioned that the calculations of § 4 on the interaction energy have been performed for several values of λ/μ and the general conclusions are much the same in all cases.

The most important physical property connected with λ/μ is the dipole moment of the molecule as a whole, so we shall get an expression for this in terms of λ/μ and then use the experimental value of 1.84×10^{-18} e.s.u. to obtain a numerical estimate.

The dipole moment of a molecule with $2N$ electrons filling N orbitals is

$$\mathbf{M} = - \int_{(2N)} \bar{\Phi} \left(\sum_{i=1}^{2N} e \mathbf{r}_i \right) \Phi (d\tau)^{2N} + \sum_{\alpha} Z_{\alpha} e \mathbf{r}_{\alpha}, \quad (3.01)$$

where Φ is the total determinantal wave function discussed in paper I, $\mathbf{r}_i$ is the position vector of the i th electron and $\mathbf{r}_{\alpha}$ that of the nucleus α . By integrating over all electrons but one for each term in the first part, this reduces to

$$\mathbf{M} = -2e \sum_{j=1}^N \int \bar{\psi}_j \psi_j \mathbf{r} d\tau + e \sum_{\alpha} Z_{\alpha} \mathbf{r}_{\alpha}. \quad (3.02)$$

Thus to evaluate the total moment we may evaluate the moments of the orbitals separately, and combine them vectorially.

To obtain an approximate value for λ/μ it is necessary to use specific analytical approximations for the atomic orbitals in equations (2.09) to (2.11). We shall use

$$\left. \begin{aligned} \phi(\text{O}; 2s) &= (Z^5/96\pi)^{1/2} r e^{-1/2 Zr}, \\ \phi(\text{O}; 2p) &= (Z^5/32\pi)^{1/2} r \cos \theta e^{-1/2 Zr}, \end{aligned} \right\} \quad (3.03)$$

for the oxygen functions with Slater's (1930) value for the screening constant Z (4.55 for oxygen). The hydrogen orbital is

$$\phi(\text{H}; 1s) = (1/\sqrt{\pi}) e^{-r_{\text{H}}}. \quad (3.04)$$

To simplify the calculations we have omitted the radial node from $\phi(\text{O}; 2s)$. This can be shown to have little effect on calculations involving only electrostatic properties of the charge distribution.

Before evaluating the dipole moment of the equivalent orbitals, we shall give a brief discussion of the charge distribution in the bond orbital. The full expression for it is

$$\chi(b_1)^2 = \lambda^2 \{\phi(\text{O}; b_1)\}^2 + 2\lambda\mu \{\phi(\text{O}; b_1) \phi(\text{H}_{(1)}; 1s)\} + \mu^2 \{\phi(\text{H}_{(1)}; 1s)\}^2. \quad (3.05)$$

In later work on the electrostatic interaction between various orbitals, the central part of (3.05) leads to some difficult integrals. The term in question represents a total charge

$$2\lambda\mu S = 2\lambda\mu \int \phi(\text{O}; b_1) \phi(\text{H}_{(1)}; 1s) d\tau$$

distributed mainly in the region between the nuclei. We shall replace it by

$$\lambda\mu S[\{\phi(\text{O}; b_1)\}^2 + \{\phi(\text{H}_{(1)}; 1s)\}^2]$$

which represents the same total charge distributed in a similar manner. As the central term in (3.05) is generally smaller than the other two, this will not introduce serious errors. The approximate form for $\chi(b_1)^2$ can now be written

$$\chi(b_1)^2 = (\lambda^2 + \lambda\mu S)\{\phi(\text{O}; b_1)\}^2 + (\lambda\mu S + \mu^2)\{\phi(\text{H}_{(1)}; 1s)\}^2. \quad (3.06)$$

It would be possible to use the exact form (3.05) for calculations on the dipole moment, but for consistency with later work, we shall use the approximation (3.06) throughout.

We can now express the total moment of the molecule in atomic units in the form

$$M = M_l + M_b + 2\rho \cos \frac{1}{2}\alpha, \quad (3.07)$$

where $M_l = \frac{40}{Z\sqrt{3}} \sin \epsilon_l \cos \epsilon_l \cos \frac{1}{2}\beta,$

$$M_b = -(\lambda^2 + \lambda\mu S) \frac{40}{Z\sqrt{3}} \sin \epsilon_b \cos \epsilon_b \cos \frac{1}{2}\alpha - 4(\lambda\mu S + \mu^2) \rho \cos \frac{1}{2}\alpha, \quad (3.08)$$

ρ being the O-H distance. M_l represents the resultant moment of the lone-pair orbitals about the O-nucleus, while M_b is that of the bond orbitals, both measured in the direction of the resultant total dipole moment. The remaining term in (3.07) is due to the two protons. Using the experimental values for ρ (0.97 Å) and α (105°) it is now possible to determine the remaining orbital parameters and consequently the resultant dipole moment as a function of λ/μ . The results are given in table 1.

TABLE 1. DIPOLE MOMENTS OF H₂O IN UNITS OF 10⁻¹⁸ E.S.U.

λ/μ	M_l	M_b	M
0.82	3.09	-7.71	0.98
1.15	3.15	-6.91	1.84
1.71	3.20	-5.85	2.95
∞	3.17	-3.17	5.60

By comparing the figures in the last column with the experimental value of 1.84×10^{-18} e.s.u., we deduce that the ratio λ/μ is approximately 1.15.

These figures are of considerable interest, as previous theories (Mulliken 1935*b*; Pauling 1940) have attributed the whole molecular dipole moment to the partial transfer of bonding electrons from the hydrogen orbitals to the more electronegative oxygen. The present theory shows, on the other hand, that a very considerable contribution arises from the orientation of the lone pairs on the side of the oxygen nucleus remote from the hydrogen nuclei.

It should be pointed out that the division of the dipole moment of the electrons into two parts M_l and M_b suffers from the same artificiality as our original definition of lone pair and bond orbitals. The only observables are the total electron distribution and the total moment. Similar difficulties prevent our defining such terms as 'the dipole moment of the O-H bond' exactly, for the contributions of lone-pair electrons can only be divided between the bonds in an artificial manner.

According to the discussion in §2, the method we have used probably slightly exaggerates the extent to which the lone-pair orbitals are oriented on the remote side of the molecule. It does appear, however, that the shift of charge from the hydrogen atoms will be considerably less than has been previously estimated.

4. FACTORS DETERMINING THE EQUILIBRIUM CONFIGURATION

As has been pointed out earlier in this paper, it is not satisfactory to discuss the equilibrium configuration of a molecule like H_2O by a method which does not allow for energy changes due to the lone pairs. The structure of the water molecule has recently been examined by Heath & Linnett (1948) who, using the experimental vibrational frequencies, show that hydrogen repulsions are not sufficient to explain the opening out of α from 90° to approximately 105° and that some hybridization must occur. The present theory allows for this, but goes further in trying to find out why such hybridization lowers the total energy and what opposing factors prevent the bond angle increasing beyond 105° . To do this we shall examine the expression for the total energy obtained in part I and endeavour to find out how some of its terms vary as the bond angle is increased from 90 to 180° , keeping the O-H distance fixed at its observed equilibrium value of $0.97 \text{ \AA}$.

The interpretation of the various parts of the expression for the total energy has been given in part IV. For the purpose of discussing the equilibrium configuration we shall divide the energy into two parts in the following way. The first part, which we shall call the local energy of the orbitals, includes the kinetic energy, energy of interaction of electrons in the same orbital and potential energy of the electrons in each equivalent orbital due to the nuclei near which that particular orbital is concentrated. The remainder, which we shall call the interaction energy, represents the interaction between different parts of the molecule. As we are chiefly interested in the effect of lone-pair interactions on the total energy variation, we shall limit ourselves to a discussion of the interaction part. In the case of H_2O the interaction energy (omitting inner shell terms and the interaction between the oxygen nucleus and the protons which remains constant if the angle only is changed) may be written

$$E_I = 1/\rho_{\text{HH}} - 4G_{bb}(x_{\text{H}}) - 8G_u(x_{\text{H}}) + 4(b_1 b_2 | g | b_1 b_2) + 4(l_1 l_2 | g | l_1 l_2) \\ + 16(bl | g | bl) - 2(b_1 b_2 | g | b_2 b_1) - 2(l_1 l_2 | g | l_2 l_1) - 8(bl | g | lb), \quad (4.01)$$

where we have written $-G_{bb}(x_{\text{H}})$ for the interaction between an electron in one bond orbital and the proton of the other bond. The first term in (4.01) represents the proton-proton repulsion, the second and third give the potential energy of electrons due to distant nuclei and the remaining terms give the interaction energy between electrons in different orbitals.

As the equivalent orbitals will be localized, the exchange terms will make a contribution only where b and l orbitals overlap, that is, well within the O-atom. The exchange terms are unlikely, therefore, to contribute effectively to the *shape* of the

molecule, so we shall neglect them. It is convenient to divide the remainder of (4.01) into three parts

$$\left. \begin{aligned} E_I(\text{bond-bond}) &= 1/\rho_{HH} - 4G_{bb}(x_H) + 4(b_1 b_2 | g | b_1 b_2), \\ E_I(\text{bond-lone pair}) &= -8G_u(x_H) + 16(bl | g | bl), \\ E_I(\text{lone pair-lone pair}) &= 4(l_1 l_2 | g | l_1 l_2). \end{aligned} \right\} \quad (4.02)$$

Of these the first represents the proton repulsion, the attraction of each bond orbital for the proton of the other and the repulsion of the orbitals; the second gives the attraction of the lone pairs by the protons and their repulsion by the bond electrons; the third corresponds to repulsion between the lone pairs. We shall investigate the variation of these three terms separately (see table 3).

Variation of orbitals with bond angle

The value $\lambda/\mu = 1.15$ deduced for the polarity of the bond orbital in § 4 strictly applies to the equilibrium configuration only, but as the results are found to be insensitive to λ/μ , we shall use this value for all bond angles. The hybridization parameters and the angles β and γ can then be obtained as functions of α from equations (2.10) and (2.13) to (2.15). The results are given in table 2.

TABLE 2. VARIATION OF ORBITALS IN H_2O WITH BOND ANGLE (DEGREES)

angle between O-H bonds α	angle between lone pairs β	γ	bond hybridization parameter ϵ_b	lone pair hybridization parameter ϵ_l
90	133.2	106.3	90	50.4
100	108.2	112.1	67.4	60.8
105	104.3	111.9	63.0	63.5
109.47	101.8	111.4	60.0	65.7
120	97.5	109.2	54.7	70.1
180	90	90	45	90

The last two columns give some indication of the degree to which hybridization occurs in both the bond orbitals (Heath & Linnett 1948) and the lone pairs. It is seen that for the equilibrium value of α (105°) both hybridization parameters are close to the value ($\epsilon = 60^\circ$) corresponding to the tetrahedral orbitals originally introduced by Pauling (1931). Also the angles α , β , γ are all of similar magnitude. Thus this equivalent orbital description of the water molecule assigns it an approximately tetrahedral character, with bond orbitals pointing in two of the directions and equivalent lone pairs in the other two. This may be regarded as a theoretical foundation for the tetrahedral model used by Bernal & Fowler (1933) and others in the discussion of the structure and properties of ice and liquid water.

Variation of interaction energy with bond angle

We shall define certain Coulomb interaction integrals between atomic functions. Suppose

$$W_{12}(\epsilon_1, \epsilon_2, \alpha_{12}) = \iint \phi_1^2(1) \phi_2^2(2) (1/r_{12}) d\tau_1 d\tau_2, \quad (4.03)$$

where ϕ_1 and ϕ_2 are atomic directed orbitals with hybridization parameters ϵ_1 and ϵ_2 (cf. equation (2.09)) on the same centre and inclined at an angle α_{12} ;

$$W_{1H}(\epsilon_1, \alpha_{1H}) = \iint \phi_1^2(1) \phi_H^2(2) (1/r_{12}) d\tau_1 d\tau_2, \quad (4.04)$$

where ϕ_H is a hydrogen orbital on another centre, whose direction is inclined at an angle α_{1H} to the directed orbital ϕ_1 ; and

$$W_{1N}(\epsilon_1, \alpha_{1N}) = \int \phi_1^2(1) (1/r_{1N}) d\tau_1, \quad (4.05)$$

where N is a hydrogen nucleus, whose direction makes an angle α_{1N} with that of ϕ_1 .

We also require two integrals involving only hydrogen functions

$$W_{HH}(\rho_{HH}) = \iint \phi_{H1}^2(1) \phi_{H2}^2(2) (1/r_{12}) d\tau_1 d\tau_2, \quad (4.06)$$

where ρ_{HH} is the separation of the two centres; and

$$W_{HN}(\rho_{HN}) = \int \phi_H^2(1) (1/r_{1N}) d\tau, \quad (4.07)$$

giving the interaction of a hydrogen orbital with a nucleus at distance ρ_{HN} . Formulae for the integrals (4.03) to (4.07) are given in the appendix.

The energy terms in (4.02) can be easily expressed in terms of these integrals. We get, using the approximate electrostatic distribution, given by (3.06)

$$\left. \begin{aligned} G_{bb}(x_H) &= (\lambda^2 + \lambda\mu S) W_{1N}(\epsilon_b, \alpha) + (\lambda\mu S + \mu^2) W_{HN}(2\rho \sin \tfrac{1}{2}\alpha), \\ (b_1 b_2 | g | b_1 b_2) &= (\lambda^2 + \lambda\mu S)^2 W_{12}(\epsilon_b, \epsilon_b, \alpha) + 2(\lambda^2 + \lambda\mu S) (\lambda\mu S + \mu^2) W_{1H}(\epsilon_b, \alpha) \\ &\quad + (\lambda\mu S + \mu^2)^2 W_{HH}(2\rho \sin \tfrac{1}{2}\alpha), \\ G_u(x_H) &= W_{1N}(\epsilon_l, \gamma), \\ (bl | g | bl) &= (\lambda^2 + \lambda\mu S) W_{12}(\epsilon_b, \epsilon_l, \gamma) + (\lambda\mu S + \mu^2) W_{1H}(\epsilon_l, \gamma), \\ (l_1 l_2 | g | l_1 l_2) &= W_{12}(\epsilon_l, \epsilon_l, \beta). \end{aligned} \right\} \quad (4.08)$$

Using the results of table 2 we can now evaluate the variation of the parts of the interaction energy (4.02). This is given in table 3.

Of the parts of the bond-bond energy, the electron repulsion and nuclear repulsion terms in equation (4.02) are partially cancelled by the attraction of the electrons for the hydrogen nucleus of the other bond. As is seen from the corresponding column of table 3, the total bond-bond interaction energy only varies slightly with bond angle. In the case of the bond-lone pair interaction energy, the variation is dominated by the electron repulsion part, so that this term has a minimum at the intermediate configuration in which the bonds and lone pairs are separated as much as possible. The lone pair-lone pair interaction energy is unmodified by any attractive term and increases as the lone pair orbitals are crushed together.

The calculated total interaction energy E_I is seen to have a minimum near $\alpha = 107^\circ$, which is fairly close to the observed equilibrium angle of 104.5° . It is also possible to estimate the bending frequency ν_2 from the variation of E_I with angle. This gives about 5000 cm.^{-1} , which is very much larger than the observed value of 1595 cm.^{-1} .

However, in view of the fact that the variable part of the energy is very small compared with the total that has to be dealt with and that many approximations have had to be made, accurate numerical results cannot be expected.

TABLE 3. VARIATION OF INTERACTION ENERGY OF H_2O WITH BOND ANGLE α

α	$G_{bb}(x_{\text{H}})$	$(b_1 b_2 g b_1 b_2)$	$G_{ll}(x_{\text{H}})$	$(bl g bl)$
90	0.4485	0.5425	0.4751	0.6452
100	0.4290	0.5310	0.4676	0.6209
105	0.4179	0.5206	0.4706	0.6191
109.47	0.4088	0.5125	0.4737	0.6203
120	0.3899	0.4938	0.4810	0.6258
180	0.3452	0.4427	0.5025	0.6509

α	E_{I} (bond-bond)	E_{I} (lone pair-lone pair)	E_{I} (bond-lone pair)	E_{I} (total)
90	0.7664	2.8416	6.5226	10.131
100	0.7683	3.0877	6.1928	10.049
105	0.7588	3.1237	6.1410	10.024
109.47	0.7526	3.1458	6.1350	10.033
120	0.7344	3.1718	6.1646	10.071
180	0.6660	3.1779	6.3941	10.238

The chief value of this calculation lies in the light it throws on the factors which determine the potential minimum. By examining the separated parts of E_{I} , we see that as α is increased from 90° , the lone pairs are forced together on the remote side, leading to an increase in their interaction energy. To begin with this is more than offset by the fact that they are removed further from the bond orbitals in the process and by the decreased bond-bond interaction. After a certain intermediate configuration, corresponding roughly to tetrahedral orientation for the four orbitals, however, the angle decreases and the bond-lone pair energy increases. It is the combination of this with the continued crushing together of the lone pairs that prevents the molecule opening into linear form.

The structure of the ammonia molecule will be determined by similar factors. The chief difference is that there is only one lone pair, so there will only be bond-bond and bond-lone pair terms in the interaction energy. These are, neglecting exchange terms as before

$$\left. \begin{aligned} E_{\text{I}}(\text{bond-bond}) &= 3/\rho_{\text{HH}} - 12G_{bb}(x_{\text{H}}) + 12(b_1 b_2 | g | b_1 b_2), \\ E_{\text{I}}(\text{bond-lone pair}) &= -6G_{ll}(x_{\text{H}}) + 12(bl | g | bl). \end{aligned} \right\} \quad (4.09)$$

The behaviour of these two parts will be similar to the corresponding terms in H_2O , the first part increasing as the bond angle increases and the second decreasing at first and then increasing. As there are no lone pair-lone pair interactions, the barrier for the plane configuration will be smaller than that preventing H_2O being linear.

The author is indebted to Professor Sir John Lennard-Jones for suggesting the problem considered in this paper and for many helpful discussions in its preparation. He is also glad to acknowledge the help of the Department of Scientific and Industrial Research.

APPENDIX. EVALUATION OF INTEGRALS

All the integrals used in this paper can be evaluated by standard methods, usually leading to simple expressions in terms of two-centre, one-electron integrals tabulated by Coulson (1942*b*). Two-centre, two-electron integrals are of the type usually described as 'Coulomb' integrals. They can be expressed in terms of

$$I_s(\kappa) = \iint r_{a1}^2 e^{-\kappa r_{a1}} P_s(\cos \theta_{1a}) e^{-r_{b2}} (1/r_{12}) d\tau_1 d\tau_2, \quad (5.01)$$

where r_{a1} is the distance from electron 1 to the centre A , θ_{1a} is the angle between this line and the axis AB and r_{b2} is the distance from electron 2 to the centre B . If the distance AB is ρ these can be expressed in terms of Coulson's J -integrals by using the subsidiary result

$$\frac{1}{\pi} \int e^{-r_{b2}} (1/r_{12}) d\tau_2 = \frac{1}{r_{b1}} - \left\{ \frac{1}{r_{b1}} + 1 \right\} e^{-r_{b1}}. \quad (5.02)$$

We then get

$$\left. \begin{aligned} \frac{1}{\pi} I_0(\kappa) &= J_5(\kappa, 0, \rho) - J_5(\kappa, 1, \rho) - J_{18}(\kappa, 1, \rho), \\ \frac{1}{\pi} I_1(\kappa) &= J_{21}(\kappa, 0, \rho) - J_{21}(\kappa, 1, \rho) - J_{20}(\kappa, 1, \rho), \\ \frac{1}{\pi} I_2(\kappa) &= \frac{3}{2} \{ J_{23}(\kappa, 0, \rho) - J_{23}(\kappa, 1, \rho) - J_{22}(\kappa, 1, \rho) \} - \frac{1}{2\pi} I_0(\kappa). \end{aligned} \right\} \quad (5.03)$$

The values of the integrals defined in the paper are summarized below, again using Coulson's J -integrals.

$$(a) \text{ Overlap integrals} \quad S_s = (Z^5/96\pi^2)^{\frac{1}{2}} J_2(\frac{1}{2}Z, 1, \rho), \quad (5.04)$$

$$S_p = (Z^5/32\pi^2)^{\frac{1}{2}} J_7(\frac{1}{2}Z, 1, \rho). \quad (5.05)$$

(b) *Electron-electron integrals*

$$\begin{aligned} Z^{-1}W_{12}(\epsilon_1, \epsilon_2, \alpha_{12}) &= \frac{93}{512} + \frac{185}{1152} \sin \epsilon_1 \cos \epsilon_1 \sin \epsilon_2 \cos \epsilon_2 \cos \alpha_{12} \\ &\quad + \frac{9}{640} \sin^2 \epsilon_1 \sin^2 \epsilon_2 P_2(\cos \alpha_{12}), \end{aligned} \quad (5.06)$$

$$\begin{aligned} W_{1H}(\epsilon_1, \alpha_{1H}) &= (Z^5/96\pi^2) I_0(Z) + (Z^5/16 \sqrt{3}\pi^2) I_1(Z) \sin \epsilon_1 \cos \epsilon_1 \cos \alpha_{1H} \\ &\quad + (Z^5/48\pi^2) I_2(Z) \sin^2 \epsilon_1 P_2(\cos \alpha_{1H}), \end{aligned} \quad (5.07)$$

$$W_{HH}(\rho_{HH}) = (1/\rho_{HH}) - \{ (1/\rho_{HH}) + \frac{11}{8} + \frac{3}{4}\rho_{HH} + \frac{1}{6}\rho_{HH}^2 \} e^{-2\rho_{HH}}. \quad (5.08)$$

(c) *Electron-nucleus interactions*

$$\begin{aligned} W_{1N}(\epsilon_1, \alpha_{1N}) &= (Z^5/96\pi) J_5(Z, 0, \rho) + (Z^5/16 \sqrt{3}\pi) J_{21}(Z, 0, \rho) \cos \epsilon_1 \sin \epsilon_1 \cos \alpha_{1N} \\ &\quad + (Z^5/32\pi) \{ J_{23}(Z, 0, \rho) - J_5(Z, 0, \rho) \} \sin^2 \epsilon_1 P_2(\cos \alpha_{1N}), \end{aligned} \quad (5.09)$$

$$W_{HN}(\rho_{HN}) = (1/\rho_{HN}) - \{ (1/\rho_{HN}) + 1 \} e^{-2\rho_{HN}}. \quad (5.10)$$

(d) *Dipole moment of an atomic directed orbital*

The electric moment of the charge distribution of an atomic directed orbital ϕ with hybridization parameter ϵ , about its centre is easily evaluated as

$$\int r_1 \cos \theta \phi^2(1) d\tau_1 = (10/Z \sqrt{3}) \sin \epsilon \cos \epsilon. \quad (5.11)$$

REFERENCES

- Bernal, J. D. & Fowler, R. H. 1933 *J. Chem. Phys.* **1**, 515.
 Coulson, C. A. 1937 *Trans. Faraday Soc.* **33**, 388.
 Coulson, C. A. 1941 *Proc. Roy. Soc. Edinb. A*, **61**, 115.
 Coulson, C. A. 1942a *Trans. Faraday Soc.* **38**, 433.
 Coulson, C. A. 1942b *Proc. Camb. Phil. Soc.* **38**, 210.
 Coulson, C. A. 1949 *J. Chim. Phys.* **46**, 198.
 Coulson, C. A., Duchesne, J. & Manneback, C. 1948 Vol. comm. Victor Henri.
 Hall, G. G. & Lennard-Jones, J. E. 1950 *Proc. Roy. Soc. A*, **202**, 155 (part III).
 Heath, D. F. & Linnett, J. W. 1948 *Trans. Faraday Soc.* **44**, 556.
 Lennard-Jones, J. E. 1929 *Trans. Faraday Soc.* **25**, 668.
 Lennard-Jones, J. E. 1949a *Proc. Roy. Soc. A*, **198**, 1 (part I).
 Lennard-Jones, J. E. 1949b *Proc. Roy. Soc. A*, **198**, 14 (part II).
 Lennard-Jones, J. E. & Pople, J. A. 1950 *Proc. Roy. Soc. A*, **202**, 166 (part IV).
 Mulliken, R. S. 1935a *J. Chem. Phys.* **3**, 375.
 Mulliken, R. S. 1935b *J. Chem. Phys.* **3**, 573.
 Pauling, L. 1931 *J. Amer. Chem. Soc.* **53**, 1367.
 Pauling, L. 1940 *Nature of the chemical bond*. Cornell University Press.
 Slater, J. C. 1930 *Phys. Rev.* **36**, 57.
 Van Vleck, J. H. & Sherman, A. 1935 *Rev. Mod. Phys.* **7**, 167.

The molecular orbital theory of chemical valency.

VI. Properties of equivalent orbitals

BY G. G. HALL, *Department of Theoretical Chemistry, University of Cambridge*

(Communicated by Sir John Lennard-Jones, F.R.S.—Received 16 December 1949)

The properties of equivalent orbitals, as defined in previous parts, are examined in more detail. It is shown that the character of an equivalent set can be deduced knowing the symmetry group of a single orbital of the set. This theorem has several useful applications. Further properties of equivalent sets are found by considering the structure of the molecular symmetry group in relation to the symmetry group of an orbital. An investigation whether or not equivalent orbitals in a molecule are uniquely defined shows that there are still degrees of freedom available which can be used to satisfy additional conditions.

1. INTRODUCTION

In parts I, II and III (Lennard-Jones 1949a, b; Hall & Lennard-Jones 1950), subsequently referred to as I, II and III, an attempt has been made to place the molecular orbital theory of valency on a more rigorous and fruitful basis. Equations determining the best possible set of one-electron orbitals in a molecule were obtained in I using the calculus of variations. The orbitals are not completely specified by these equations, so that in II several types of orbital were introduced. These different orbitals can be transformed into one another without changing the value of the determinantal wave function for the molecule, yet they each lead to certain simplifications in the equations, which are of physical significance. In III three types of orbital were distinguished, molecular orbitals, symmetry orbitals and equivalent

orbitals. *Molecular orbitals* were defined as orbitals for which a matrix of parameters E_{mn} occurring in the equations is diagonal. Orbitals belonging to one of the irreducible representations of the symmetry group of the molecule were called *symmetry orbitals*, and molecular orbitals were shown to be a particular example of them. On the other hand, a set of *equivalent orbitals* was defined by the property that the members of the set were identical with each other, except for orientation and position in space. The particular significance of molecular orbitals was examined and it was found that, when a change of energy is involved, as in spectroscopic properties, molecular orbitals provide the most suitable description of the molecule.

The object of this paper is to study the symmetry properties of equivalent orbitals and to discuss their chemical significance. Several theorems are proved concerning the symmetry of a single equivalent orbital and the definition of these orbitals in a molecule. This leads to a discussion of the structure of saturated molecules using several completely occupied sets of equivalent orbitals.

2. THE CHARACTER OF A SET OF EQUIVALENT ORBITALS

In order to apply the methods of group theory to equivalent orbitals it is necessary to re-state the definition given above in terms of a permutation group $\mathcal{G}$. This group is normally the symmetry group of the molecule, but it may occasionally be more convenient to use a subgroup of this group. A set of equivalent orbitals can then be defined by the two properties of being composed of equivalent orbitals and of being a set. Orbitals are equivalent if any member of the set can be transformed into any other member by some permutation operation of the group $\mathcal{G}$. These orbitals will form a set if any operation of $\mathcal{G}$ permutes the members of the set among themselves.

The second of these properties ensures that a set of equivalent orbitals as a whole forms a basis for a representation of the group $\mathcal{G}$ in which the various group operations are represented by permutation matrices. If the operations of the group are denoted by $E, A, B, \dots, G, \dots, P$, then they will be represented on this basis by permutation matrices $E_{ij}, A_{ij}, \dots, G_{ij}, \dots, P_{ij}$. The irreducible representations into which the representation can be separated are determined by the character of the set. The character of any equivalent set is given by the theorem that the component $\chi(\pi)$ of the character of a set of equivalent orbitals, in the class π , is a positive integer

$$\chi(\pi) = ns/p, \quad (2.01)$$

where n is the number of orbitals in the set, s is the number of operations in the class π under which any one orbital is invariant and p is the total number of operations in the class.

To prove this theorem we consider classes of operations rather than the operations themselves. These may be written formally (Dirac 1947) as

$$\epsilon = E; \quad \alpha = 1/a(A + B + \dots); \quad \dots; \quad \pi = 1/p(P + Q + \dots). \quad (2.02)$$

The corresponding matrix equations are

$$\epsilon_{ij} = E_{ij}; \quad \alpha_{ij} = 1/a(A_{ij} + B_{ij} + \dots); \quad \dots; \quad \pi_{ij} = 1/p(P_{ij} + Q_{ij} + \dots). \quad (2.03)$$

These classes have the property of being invariant under any element G of the group

$$\pi = G\pi G^{-1}, \quad (2.04)$$

and therefore

$$\pi_{ij} = \sum_{lm} G_{il} \pi_{lm} G_{mj}^{-1}. \quad (2.05)$$

The effect of pre-multiplying a matrix by a permutation matrix G_{il} is to permute the order of its rows, while post-multiplication by the inverse matrix G_{mj}^{-1} , which is also the transpose of G_{jm} , applies the same permutation to its columns. Consequently, when both permutations are applied, as in (2.05), the diagonal elements π_{ll} are permuted among themselves alone. Since G can be any element of $\mathcal{G}$ and any orbital is transformed into any other one by some element of $\mathcal{G}$, equation (2.05) implies that all the diagonal elements of π_{ij} are equal. Each of these diagonal elements must be a positive integer r divided by p , for they are defined by (2.03), and the elements of a permutation matrix are unity or zero. The trace of π_{ij} , which equals the component of the character in π , is, therefore,

$$\chi(\pi) = \text{tr}(\pi_{ij}) = nr/p. \quad (2.06)$$

Unity in the diagonal of a permutation matrix P_{ij} means that the corresponding orbital is invariant under the element P of the group $\mathcal{G}$. So this definition of r implies that the total number of such orbitals for all the elements in the class π is nr . Since the n equivalent orbitals have the same relative symmetry, r must therefore be equal to the number of elements in π under which each orbital is invariant

$$r = s. \quad (2.07)$$

Finally, by the well-known property of a class

$$\text{tr}(\pi_{ij}) = \text{tr}(P_{ij}) = \text{tr}(Q_{ij}) = \dots, \quad (2.08)$$

and each of these, and hence $\chi(\pi)$, must be a positive integer.

3. APPLICATIONS TO MOLECULAR STRUCTURE

The theorem proved in the previous section has several important uses, some of which are discussed in this section. First, if the character of an equivalent set is known, the symmetry of an individual orbital in the set can be deduced. Methane may be taken as an example. The molecular orbital of lowest energy belongs to the representation A_1 of the group T_d and is approximately the same as the carbon $1s$ orbital. This orbital, or, indeed, one which is totally symmetric under any group, can be considered as an equivalent set composed of one orbital. The next lowest molecular orbitals are of symmetry A_1 and T_2 . From these, four equivalent orbitals can be obtained (II) so that for this case $n = 4$. The symmetry of one of these equivalent orbitals can be deduced in the form of a table:

T_d	E	C_3	C_2	σ_d	S_4
number of elements in class = p	1	8	3	6	6
character of $(A_1 + T_2) = \chi$	4	1	0	2	0
$s = \frac{1}{4}p\chi$	1	2	0	3	0

This shows that the four equivalent orbitals into which the molecular orbitals spanning the A_1 and T_2 representations can be transformed have the symmetry elements $E, 2C_3, 3\sigma_d$. Thus each orbital has the full symmetry of a C_{3v} subgroup of T_d and the equations for the orbitals must have the same symmetry.

The theorem can also be used as a necessary condition that a number of symmetry or molecular orbitals can be transformed into an equivalent set. Since, for equivalent sets, the only quantities that can vary in equation (2.01) are the numbers s , it is necessary that each component of the total character of the given orbitals should satisfy (2.01) for some positive integer s . This may be illustrated by the example of the occupied π orbitals for the ground state of benzene, which cannot be transformed into an equivalent set. The three occupied π orbitals belong to the representations A_1 and E_1 . For these π orbitals the symmetry group D_{6h} of the molecule cannot be used as the permutation group, for the operation σ_h does not leave the set invariant. The subgroup C_{6v} can, however, be used and the table drawn up as before:

C_{6v}	E	C_6	C_3	C_2	σ_v	σ_d
p	1	2	2	1	3	3
character of $(A_1 + E_1)$	3	2	0	-1	1	1
$s = \frac{1}{3}p\chi$	1	$1\frac{1}{3}$	0	$-\frac{1}{3}$	1	1

In this example the theorem proves clearly that the three occupied orbitals cannot be transformed into equivalent orbitals, for neither fractional nor negative values of s are possible. It is not true, however, that orbitals can be transformed into an equivalent set if the condition (2.01) alone is satisfied. A sufficient condition for an equivalent set is obtained if the symmetry elements indicated by equation (2.01) form a subgroup $\mathcal{H}$ of the symmetry group of the molecule. This condition is obviously necessary, since the symmetry elements of any distribution in space must form a group.

4. AN ISOMORPHISM BETWEEN EQUIVALENT ORBITALS AND COSETS

We have seen that the character of an equivalent set can be used to find the subgroup $\mathcal{H}$ of $\mathcal{G}$, which is the symmetry group of one particular orbital. Knowing this subgroup we can now deduce further properties of equivalent orbitals.

This subgroup $\mathcal{H}$ determines another permutation representation of the group $\mathcal{G}$. The basis for this representation is the set of left cosets of $\mathcal{H}$ in $\mathcal{G}$. A left coset $\mathcal{H}_A$ of $\mathcal{H}$ in $\mathcal{G}$ is defined as the collection of elements of $\mathcal{G}$ which can be expressed in the form AX , where A is a fixed element of $\mathcal{G}$ and X is some element of $\mathcal{H}$. By suitable choice of the elements $E, A, \dots, P$, a set of left cosets can be found which contains each element of $\mathcal{G}$ once and once only,

$$\mathcal{H}_E = \mathcal{H}; \quad \mathcal{H}_A = A\mathcal{H}; \quad \dots; \quad \mathcal{H}_P = P\mathcal{H}. \quad (4.01)$$

Under left multiplication by any element G of $\mathcal{G}$ these cosets are permuted among themselves:

$$G\mathcal{H}_E = \mathcal{H}_G; \quad G\mathcal{H}_A = \mathcal{H}_{GA}; \quad \dots; \quad G\mathcal{H}_P = \mathcal{H}_{GP}; \quad (4.02)$$

so that the element G is represented by a permutation matrix. The order of this matrix is the index d of $\mathcal{H}$ in $\mathcal{G}$, which is

$$d = g/h, \quad (4.03)$$

where g and h are the orders of $\mathcal{G}$ and $\mathcal{H}$ respectively. Representations of this type have been studied by Littlewood (1935) in relation to the structure of abstract groups.

We shall now prove that this coset representation is the same as the equivalent orbital representation. This follows from the fact that the cosets and the equivalent orbitals are isomorphic. In this isomorphism the coset $\mathcal{H}_E$ corresponds to the particular orbital χ_1 which has $\mathcal{H}$ as its symmetry group, because any element H of $\mathcal{H}$ leaves both coset and orbital invariant

$$H\mathcal{H}_E = HE\mathcal{H} = \mathcal{H}_E; \quad H\chi_1 = \chi_1. \quad (4.04)$$

Again, an element A of $\mathcal{G}$, which permutes the orbital χ_1 into another orbital χ_2 , permutes $\mathcal{H}_E$ into $\mathcal{H}_A$

$$A\chi_1 = \chi_2; \quad A\mathcal{H}_E = \mathcal{H}_A, \quad (4.05)$$

and this coset $\mathcal{H}_A$ is isomorphic with the second orbital. The symmetry group of the coset $\mathcal{H}_A$ is $A\mathcal{H}A^{-1}$, the transform of $\mathcal{H}$ by A , because if F belongs to this group,

$$F = AXA^{-1}, \quad (4.06)$$

for some X in $\mathcal{H}$, and

$$F\mathcal{H}_A = FA\mathcal{H} = AXA^{-1}A\mathcal{H} = AX\mathcal{H} = \mathcal{H}_A. \quad (4.07)$$

This group $A\mathcal{H}A^{-1}$ is also the symmetry group of χ_2 , for

$$F\chi_2 = FA\chi_1 = AXA^{-1}A\chi_1 = AX\chi_1 = A\chi_1 = \chi_2, \quad (4.08)$$

so that the isomorphism is proved. In this way, it follows that the set of cosets is isomorphic with the set of equivalent orbitals and the two permutation representations are the same.

Several properties of equivalent sets can be deduced immediately from this theorem. In the first place, the number of orbitals in an equivalent set is equal to the number of cosets and, hence, to the index of $\mathcal{H}$ in $\mathcal{G}$:

$$n = d = g/h. \quad (4.09)$$

Thus the character of an equivalent set is entirely determined by the symmetry group $\mathcal{H}$ of one of its orbitals.

Secondly, the number of possible types of equivalent orbitals under a group $\mathcal{G}$ can be found. Any subgroup of $\mathcal{G}$ can be used to form a permutation representation and, if the subgroups are not equivalent to each other, these representations will be distinct. The number of different possible equivalent sets is, therefore, equal to the number of non-equivalent subgroups of $\mathcal{G}$. In the group T_d , for example, there are eleven such subgroups and hence eleven different types of equivalent orbital.

The number of times the totally symmetric representation occurs in the permutation representation generated by an equivalent set can also be found. It is given by the usual formula

$$\sum_{\pi} p\chi(\pi)/g. \quad (4.10)$$

Using equations (2.01) and (4.09) together with the fact that $\sum_{\pi} s$ is the number of elements in $\mathcal{H}$, that is, h ; this becomes

$$\sum_{\pi} p\chi(\pi)/g = \sum_{\pi} sn/g = \sum_{\pi} s/h = 1. \quad (4.11)$$

Among the symmetry or molecular orbitals which form a basis for an equivalent set there will, therefore, be one and only one which is totally symmetric. It is readily seen that this orbital equals the sum of the equivalent orbitals suitably normalized

$$\psi = (1/\sqrt{n})(\chi_1 + \chi_2 + \dots + \chi_n). \quad (4.12)$$

This result is of considerable importance in the theory of molecular structure for, if equivalent sets can be formed, the number of sets in a molecule is equal to the number of occupied molecular orbitals which are totally symmetric.

5. DEGREES OF FREEDOM IN EQUIVALENT ORBITALS

When a molecule contains several equivalent sets then, as we have just seen, there will be several molecular orbitals belonging to the same irreducible representation. In IV (Lennard-Jones & Pople 1950) several examples of this are discussed, and it is shown that the equivalent orbitals are not then uniquely defined by the equivalence condition alone. In this section the definition of equivalent orbitals for any molecule containing several equivalent sets will be considered. As a start, equivalent orbitals are defined as transforms of a set of known orbitals and the properties of the irreducible representations used to show that this definition is not unique.

Since molecular orbitals belong to the various irreducible representations of the group, the effect of an element G of the group on the set of occupied molecular orbitals is represented by equations of the form

$$G\psi_i = \sum_j \Gamma_{ij} \psi_j. \quad (5.01)$$

This matrix Γ_{ij} has block form; that is, its elements are zero except for square blocks of terms around the diagonal, each block corresponding to one particular occurrence of an irreducible representation. Thus in the ammonia molecule, for example (IV), each matrix Γ_{ij} has four blocks of terms on its diagonal. The first three are of dimensions one by one and correspond to the three occupied A_1 orbitals, while the fourth has dimensions two by two corresponding to the two orbitals of the twofold degenerate E representation. From these molecular orbitals equivalent orbitals may be defined, as in equations (2.06) and (2.07) of III, by

$$\chi_i = \sum_j T_{ij} \psi_j, \quad (5.02)$$

$$\psi_i = \sum_j T_{ij}^{-1} \chi_j = \sum_j \bar{T}_{ji} \chi_j. \quad (5.03)$$

The equivalent orbitals are affected by G in the same way as by a permutation matrix

$$G\chi_i = \sum_j G_{ij} \chi_j, \quad (5.04)$$

but now G_{ij} will also have block form with a block corresponding to each equivalent set. In ammonia there will be three such blocks. Two of these are of dimensions one by one and correspond to the inner shell and the lone pair associated with the nitrogen atom. The third has dimensions three by three and corresponds to the

representation set up by the three NH bond equivalent orbitals. The equation connecting these matrices has been deduced in III and is

$$G_{kl} = \sum_{ij} T_{ki} \Gamma_{ij} T_{jl}^{-1}. \quad (5.05)$$

We now assume that the molecular orbitals ψ_i and hence the matrices Γ_{ij} are uniquely fixed. The permutation matrices are also unique once the order of the orbitals χ_i has been fixed, but these orbitals themselves need not be unique. If, therefore, there exists another set of equivalent orbitals χ'_n with the same properties and defined by

$$\chi'_n = \sum_m S_{nm} \psi_m, \quad (5.06)$$

then, by the same argument as before,

$$G_{kl} = \sum_{mn} S_{km} \Gamma_{mn} S_{nl}^{-1}. \quad (5.07)$$

By the uniqueness of the permutation matrix, equations (5.05) and (5.07) can be combined, so that

$$\sum_{ij} T_{ki} \Gamma_{ij} T_{jl}^{-1} = \sum_{mn} S_{km} \Gamma_{mn} S_{nl}^{-1}, \quad (5.08)$$

and this leads to

$$\Gamma_{il} = \sum_{klmn} T_{ik}^{-1} S_{km} \Gamma_{mn} S_{nl}^{-1} T_{lj}. \quad (5.09)$$

Defining a matrix R_{im} by the equation

$$R_{im} = \sum_k T_{ik}^{-1} S_{km} \quad (5.10)$$

enables (5.09) to be written as

$$\Gamma_{ij} = \sum_{mn} R_{im} \Gamma_{mn} R_{nj}^{-1}, \quad (5.11)$$

or

$$\sum_j \Gamma_{ij} R_{jn} = \sum_m R_{im} \Gamma_{mn}. \quad (5.12)$$

It is convenient to partition these matrices so that the blocks of the group matrices are in separate partitions. Equation (5.12) can then be written

$$\sum_J \Gamma_{IJ} R_{JN} = \sum_M R_{IM} \Gamma_{MN}, \quad (5.13)$$

where each element of these matrices is itself a matrix. The group matrices, however, have non-zero terms in blocks along the diagonal only, so they can be written

$$\Gamma_{MN} = \Gamma_M \delta_{MN} \quad (5.14)$$

and (5.13) becomes

$$\sum_J \Gamma_I \delta_{IJ} R_{JN} = \sum_M R_{IM} \Gamma_M \delta_{MN}, \quad (5.15)$$

$$\Gamma_I R_{IN} = R_{IN} \Gamma_N. \quad (5.16)$$

If I and N refer to *different* irreducible representations Γ_I and Γ_N of the group element G , then the conditions of Schur's lemma are satisfied (see, for example, Weyl 1931), and it follows that

$$R_{IN} = 0. \quad (5.17)$$

If $I = N$ this lemma can again be applied and, since R is unitary,

$$R_{NN} = \pm I. \quad (5.18)$$

On the other hand, if I and N refer to the same irreducible representation there are no restrictions on the matrix $R_{I,N}$ so that the matrix T_{ij} which defines the equivalent orbitals will have degrees of freedom corresponding to arbitrary unitary transformations of orbitals belonging to the same irreducible representation.

This result may be expressed more simply by saying that the sets of equivalent orbitals for a molecule are formed from symmetry orbitals belonging to fixed irreducible representations of the group, but that the particular orbitals used when there are more than one of the same symmetry are not specified. The degrees of freedom introduced into the orbitals by means of the transformation T_{ij} will only be fully utilized when the correct linear combination of symmetry orbitals of the same representation has been decided in each case by some additional criterion. Thus in IV it is found satisfactory to use the lowest molecular orbital of ammonia as the inner shell equivalent orbital; but, in forming the lone pair and NH bond equivalent orbitals, the orbitals are more localized if symmetry orbitals are constructed as linear combinations of the two remaining A_1 molecular orbitals.

6. THE STRUCTURE OF MOLECULES

From the practical point of view there are two advantages in using equivalent orbitals to investigate molecular structure. Since the members of a set are equivalent to each other it is only necessary to solve one equation to determine the whole set. Then again, since each equivalent orbital must belong to the totally symmetric representation of its own symmetry subgroup, the equations must have definite symmetry properties. The importance of equivalent orbitals does not depend, however, on these practical advantages but rather on their chemical significance.

The significance of equivalent orbitals has been discussed by Lennard-Jones & Pople (1950) in IV. They have shown that the 'exchange' contribution to the total energy is small when equivalent orbitals are used, so that it is a good approximation to consider the electrons in equivalent orbitals as charge distributions interacting according to the Coulomb law of force. Using this result, they have also shown that these electrons may be regarded as localized around one or two nuclei. Such electrons are described as belonging to inner shells, lone pairs, or valence bonds; and the corresponding equivalent orbitals are approximately equal to atomic orbitals, directed atomic orbitals, and diatomic orbitals respectively. These results may also be obtained in a less rigorous manner by using 'linear combinations of atomic orbitals' as approximate molecular orbitals. The transformation to equivalent orbitals then shows that these orbitals reduce to one or more atomic orbitals as before. Seitz (1940) has applied this result to the inner shells of metals. A more accurate approximation to these equivalent orbitals involves, however, small contributions from the orbitals of neighbouring atoms, so that the localization is only relative. Thus in methane the CH bond equivalent orbital involves mainly a carbon tetrahedral orbital and a hydrogen orbital but has also a smaller contribution involving the other three hydrogen atoms equally (Coulson 1942).

The resemblances between this equivalent orbital description of a molecule and that given by a simple chemical approach suggest another application of the theory.

One of the difficulties in applying the theory of molecular orbitals is that it is difficult, except possibly for central molecules, to find out which molecular orbitals are occupied in the ground state. If we make the hypothesis that the electrons in the ground state of saturated molecules exactly fill a number of equivalent sets, each set corresponding to inner shell, lone pair, or valence bond electrons in the usual chemical sense, then we have a new basis for understanding the structure of these molecules. Thus from the chemical structure the characters of the corresponding equivalent sets can be found and hence, by reduction into irreducible representations, the occupied molecular orbitals. Since the only chemical property used is the symmetry of the equivalent orbitals, this hypothesis does not imply a strict localization of the electrons.

The author wishes to acknowledge his thanks to Professor Sir John Lennard-Jones for much helpful criticism and advice and to the Ministry of Education, Northern Ireland, for a Research grant.

REFERENCES

- Coulson, C. A. 1942 *Trans. Faraday Soc.* **38**, 433.
Dirac, P. A. M. 1947 *The principles of quantum mechanics*, 3rd ed. p. 214. Oxford University Press.
Hall, G. G. & Lennard-Jones, J. E. 1950 *Proc. Roy. Soc. A*, **202**, 155 (part III).
Lennard-Jones, J. E. 1949*a* *Proc. Roy. Soc. A*, **198**, 1 (part I).
Lennard-Jones, J. E. 1949*b* *Proc. Roy. Soc. A*, **198**, 14 (part II).
Lennard-Jones, J. E. & Pople, J. A. 1950 *Proc. Roy. Soc. A*, **202**, 166 (part IV).
Littlewood, D. E. 1935 *Proc. Lond. Math. Soc.* (2), **39**, 150.
Seitz, F. 1940 *Modern theory of solids*, p. 301. New York: McGraw-Hill.
Weyl, H. 1931 *The theory of groups and quantum mechanics*, p. 153. London: Methuen.

Finite strains in an anisotropic elastic continuum

BY J. G. OLDROYD

Courtaulds Limited, Research Laboratory, Maidenhead, Berks.

(Communicated by A. H. Wilson, F.R.S.—Received 2 February 1950)

The discussion in a previous paper (Oldroyd 1950), on the invariance properties required of the equations of state of a homogeneous continuum, is extended by taking into account thermodynamic restrictions on the form of the equations, in the case of an elastic solid deformed from an unstressed equilibrium configuration. The general form of the finite strain-stress-temperature relations, expressed in terms of a free-energy function, is deduced without assuming that the material is isotropic. The results of other authors based on the assumption of isotropy are shown to follow as particular cases. The equations of state are derived by considering quasi-static changes in an elastic solid continuum; the results then apply to non-ideally elastic solids in equilibrium, or subjected to quasi-static changes only, and to ideally elastic solids in general motion. A necessary and sufficient compatibility condition for the finite strains at different points of a continuum is also derived. As a simple illustration of the derivation and use of equations of state involving anisotropic physical constants, the torsion of an anisotropic cylinder is discussed briefly.

1. INTRODUCTION

The invariance properties required of rheological equations of state of general validity have been discussed in a previous paper (Oldroyd 1950), assuming only homogeneity and continuity of the material considered. The restrictions imposed on the form of the equations by the fact that rheological properties are independent of the frame of reference, and independent of the motion of the material in space, were taken into account, but no consideration of the further restrictions which must be imposed by the laws of thermodynamics was attempted in the general case. In the present paper, the conclusions of the earlier paper are applied to the case of an elastic solid deformed to a finite extent from an initial unstressed equilibrium configuration, making use of thermodynamic principles. For this purpose, an elastic continuum is considered, in a state of motion which is sufficiently slow to permit the changes in an arbitrary element to be treated as quasi-static, but which is otherwise arbitrary. It is not necessary to distinguish between isotropic and anisotropic solids in developing the theory in this way.

It is convenient to make the following distinctions in terminology. The term *elastic* will be used to cover any solid with the following property: If a material element is subjected to an arbitrary finite constant deformation and maintained at a constant temperature, the stresses (eventually) attain constant values depending only on the deformation and on the temperature, and not in any way on the previous strain history of the element. If the final stresses are attained instantaneously, the solid is called *ideally elastic*; if only after a finite or infinite time, *non-ideally elastic*. Only in the case of ideal elasticity will the stress-deformation-temperature relations be independent of the rate at which changes take place. The results concerning the form of the equations of state, obtained by considering quasi-static

changes, will therefore be applicable to ideally elastic solids in general motion, and to non-ideally elastic solids when in equilibrium or when subjected to quasi-static changes only.

The three-dimensional tensor notation and the terminology of the previous paper are retained. Tensors may be referred either to a system of space co-ordinates ξ^i convected with the continuum under consideration, or to a system of co-ordinates x^i fixed in space. The convected- and fixed-components of the same tensor field associated with the material at an instant t are then denoted by corresponding small Greek and Latin letters (except that no such distinction is necessary for absolute scalars), e.g. $\beta_{:j}^{:l}(\xi, t)$, $b_{:i}^{:k}(x, t)$. The corresponding capital Latin letter may then be used to denote Eulerian fixed-components; thus, $B_{:i}^{:k}(x, t, t_0)$ is the tensor at x^i with the same convected-components ($\beta_{:j}^{:l}(\xi, t_0)$ at ξ^j) as the tensor $b_{:i}^{:k}(x_0, t_0)$, primarily associated with the material (x^i, t) when it was in the position x_0^i at the previous instant t_0 . The usual summation convention applies to tensor suffixes, i.e. to small Latin suffixes.

2. THE EQUATIONS OF STATE

By making use of the general conclusions of Oldroyd (1950), it is possible to write down the quantities which will be involved in the invariant equations of state of an elastic solid. A convected system of co-ordinates ξ^i is first set up in an elastic solid continuum moving in an arbitrary manner, and the convected-components of the metric tensor at time t are denoted by $\gamma_{jl}(\xi, t)$. At a particular instant t_0 , the material is assumed to be in equilibrium with vanishing stresses everywhere. To fix ideas, it is supposed that the temperature T has the constant value T_0 at time t_0 . By the definition of ideal elasticity, the stress tensor† $\pi_{jl}(\xi, t)$ at ξ^j in an ideally elastic solid at time t is defined completely by the instantaneous configuration and temperature. Therefore the equations of state are essentially invariant *algebraic* relations, of which six are independent, between the quantities

$$\gamma_{jl}(\xi, t), \quad \gamma_{jl}(\xi, t_0), \quad \pi_{jl}(\xi, t), \quad T(\xi, t), \quad T_0, \quad \kappa_N{}_{:j}^{:l}(\xi) \quad (N = 1, 2, 3, \dots), \quad (1)$$

where the tensors $\kappa_N{}_{:j}^{:l}$ are physical constants of the material. If the material is elastic but not ideally elastic, the equations of state will in general take the form of integro-differential equations relating the following functions of the time:

$$\gamma_{jl}(\xi, t'), \quad \pi_{jl}(\xi, t'), \quad T(\xi, t'), \quad \kappa_N{}_{:j}^{:l}(\xi) \quad (t_0 \leq t' \leq t; N = 1, 2, 3, \dots). \quad (2)$$

But these must reduce to algebraic relations between the quantities (1) when an equilibrium state is assumed.

In transforming the equations of state to a fixed system of co-ordinates x^i , every tensor which occurs in the equations is represented by fixed-components at a single point x^i , the position in space of the material ξ^j at the current time t . The metric and stress tensors at x^i have fixed-components $g_{ik}(x)$ and $p_{ik}(x, t)$. The initial-metric

* ξ is used as an abbreviation for ξ^1, ξ^2, ξ^3 ; x for x^1, x^2, x^3 .

† The symmetric second-order absolute tensor is indicated. Any other tensor measure of stress is expressible in terms of this stress tensor and the metric tensor.

tensor (with convected-components $\gamma_{jl}(\xi, t_0)$) is represented by the fixed-components $G_{ik}(x, t, t_0)$, which can be expressed in terms of the known fixed-components of the metric tensor field and the displacement functions

$$x_0^i = x_0^i(x, t, t_0) \quad (3)$$

$$\text{as} \quad G_{ik}(x, t, t_0) = g_{mp}(x_0) \frac{\partial x_0^m}{\partial x^i} \frac{\partial x_0^p}{\partial x^k} \quad (4)$$

The physical constants have fixed-components

$$k_{N:i:k:\dots}(x, t) = K_{N:i:k:\dots}(x, t, t_0) = \left\| \frac{\partial x_0}{\partial x} \right\|^W \Pi \left(\frac{\partial x_0^m}{\partial x^i} \right) \Pi' \left(\frac{\partial x^k}{\partial x_0^p} \right) k_{N:\dot{m}:p:\dots}(x_0, t_0), \quad (5)$$

where Π (Π') denotes a product of similar terms, one for each covariant (contravariant) suffix, W is the weight of the tensor and $\| \partial x_0 / \partial x \|$ denotes the Jacobian determinant of the $\partial x_0^k / \partial x^i$'s; the components $k_{N:\dot{m}:p:\dots}(x_0, t_0)$ of the physical constants associated with the material in the undeformed state may usually be supposed known functions of position. It will not be necessary in the present paper to consider Eulerian components of tensors primarily associated with instants $t' \neq t_0$, and we may without ambiguity drop the explicit mention of t_0 in $G_{ik}(x, t, t_0)$, $K_{N:i:k:\dots}(x, t, t_0)$, etc., and also in the displacement functions (3). The transformed equations of state of an ideally elastic solid, or of a non-ideally elastic solid in equilibrium, may then be written as a set of algebraic relations between the following tensors at x^i :

$$g_{ik}(x), \quad G_{ik}(x, t), \quad p_{ik}(x, t), \quad T(x, t), \quad T_0, \quad K_{N:i:k:\dots}(x, t) \quad (N = 1, 2, 3, \dots). \quad (6)$$

If the motion of an elastic continuum is sufficiently slow for the changes in an arbitrary element to be treated as quasi-static, the thermodynamic state of the material at x^i at time t is completely defined by the variables (6), from which p_{ik} can be omitted without loss of generality. It is known that the instantaneous rate of working of the external forces on an element of any continuum exceeds the rate of increase of kinetic energy by an amount $p^{ik}e_{ik}^{(1)}$ per unit volume, where $e_{ik}^{(1)} (\equiv \frac{1}{2} \delta g_{ik} / \delta t)$ is the rate-of-strain tensor. The fundamental thermodynamic formula relating the free-energy per unit mass F with the tensors

$$g_{ik}, \quad G_{ik}, \quad T, \quad T_0, \quad K_{N:i:k:\dots} \quad (N = 1, 2, 3, \dots) \quad (7)$$

$$\text{therefore takes the form} \quad \frac{\delta F}{\delta t} = \frac{p^{ik}}{2\rho} \frac{\delta g_{ik}}{\delta t} - S \frac{\delta T}{\delta t}, \quad (8)$$

where ρ is the density, and $\delta / \delta t$ denotes convected differentiation in relation to the continuum in an arbitrary quasi-static motion. (This is of the expected form, since G_{ik} , T_0 and the K_N 's all vanish on convected differentiation. S is the entropy per unit mass and $\frac{1}{2} p^{ik} / \rho$ can be regarded as the generalized component of force corresponding to the generalized co-ordinate g_{ik} .) Allowing for the symmetry of $\delta g_{ik} / \delta t$ and p^{ik} , we have

$$p^{ik} = \rho \Delta^{ik} F, \quad (9)$$

where

$$\Delta^{ik} \equiv \frac{\partial}{\partial g_{ik}} + \frac{\partial}{\partial g_{ki}}. \quad (10)^*$$

* Whenever this abbreviation is used, the differentiation with respect to g_{ik} is understood to be performed holding all other components of the tensors (7) constant (including g_{ki}).

Equation (9) gives the general form of the equations of state for an ideally elastic solid in general motion, or for a non-ideally elastic solid in equilibrium. There is a restriction on the possible forms of F , regarded as an absolute scalar combination of the tensors (7), derivable from the vanishing of the stresses in the undeformed material; $\Delta^{ik}F$ must vanish when $G_{ik} = g_{ik}$ and $T = T_0$.

In equation (9), the density ρ may be expressed in terms of the equilibrium density ρ_0 at temperature T_0 and the determinants g and G of the g_{ik} 's and G_{ik} 's as

$$\rho = \rho_0 (G/g)^{\frac{1}{2}}; \quad (11)$$

this is a form of the equation of continuity.

3. ALTERNATIVE FORMS OF THE EQUATIONS

The extremely simple form of the equations of state (9) can only be used directly if a given expression for F is first reduced to a combination of the tensors (7); but the reduction to this form may in practice be laborious. The function F may be given as an absolute scalar combination of tensors

$$h_{M:i:k\cdots}, \quad T, \quad T_0, \quad K'_N{}_{:i:k\cdots} \quad (M, N = 1, 2, 3, \dots), \quad (12)$$

where the K'_N 's are a set of physical constants (which will be combinations of the K_N 's of (7) and universally constant tensors) and the h_M 's are a set of tensors measuring deformation, not necessarily independent, but each reducible to a combination of the tensors g_{ik} , G_{ik} and universal constants. The generalized form of (9) is clearly

$$p^{ik} = \rho \sum_M \frac{\partial F}{\partial h_{M:m:p\cdots}} \Delta^{ik} h_{M:m:p\cdots}. \quad (13)$$

It will often happen that F can be expressed without difficulty in terms of G_{ik} , the symmetric reciprocal tensor $\bar{G}^{ik}$ defined by

$$G_{im} \bar{G}^{km} = \delta_i^k, \quad (14)$$

tensors obtained from these by raising or lowering suffixes, and invariants obtained by contraction. The notation

$$J_1 \equiv \delta_k^i G_i^k, \quad J_2 \equiv \frac{1}{2} \delta_{kp}^{im} G_i^k G_m^p, \quad J_3 \equiv \frac{1}{6} \delta_{kpt}^{imr} G_i^k G_m^p G_r^t \quad (15)$$

will be used for the invariants of G_i^k , and I_1, I_2, I_3 denote the corresponding invariants of $\bar{G}_i^k$.* For later reference, it is convenient to record the following formulae for use in conjunction with (13):

$$\Delta^{ik} g_{mp} = \delta_m^i \delta_p^k + \delta_m^k \delta_p^i, \quad \Delta^{ik} g^{mp} = -g^{im} g^{kp} - g^{km} g^{ip}, \quad (16)$$

$$\Delta^{ik} G_{mp} = 0, \quad \Delta^{ik} G_m^p = -g^{ip} G_m^k - g^{kp} G_m^i, \\ \Delta^{ik} G^{mp} = -g^{im} G^{kp} - g^{ip} G^{km} - g^{km} G^{ip} - g^{kp} G^{im}, \quad (17)$$

$$\Delta^{ik} \bar{G}^{mp} = 0, \quad \Delta^{ik} \bar{G}_m^p = \delta_m^i \bar{G}^{kp} + \delta_m^k \bar{G}^{ip}, \quad \Delta^{ik} \bar{G}_{mp} = \delta_m^i \bar{G}_p^k + \delta_p^i \bar{G}_m^k + \delta_m^k \bar{G}_p^i + \delta_p^k \bar{G}_m^i, \quad (18)$$

$$\Delta^{ik} J_1 = -2G^{ik}, \quad \Delta^{ik} J_2 = 2(G^{im} G_m^k - J_1 G^{ik}) = 2(J_3 \bar{G}^{ik} - J_2 g^{ik}), \quad \Delta^{ik} J_3 = -2J_3 g^{ik}, \quad (19)$$

$$\Delta^{ik} I_1 = 2\bar{G}^{ik}, \quad \Delta^{ik} I_2 = 2(I_1 \bar{G}^{ik} - \bar{G}^{im} \bar{G}_m^k) = 2(I_2 g^{ik} - I_3 G^{ik}), \quad \Delta^{ik} I_3 = 2I_3 g^{ik}. \quad (20)$$

* In terms of the principal extension ratios usually denoted by $\lambda_1, \lambda_2, \lambda_3$, $I_1 = \lambda_1^2 + \lambda_2^2 + \lambda_3^2$, $I_2 = \lambda_1^2 \lambda_2^2 + \lambda_2^2 \lambda_3^2 + \lambda_1^2 \lambda_3^2$, $I_3 = \lambda_1^2 \lambda_2^2 \lambda_3^2$; J_1, J_2, J_3 are the corresponding expressions with λ_α replaced by $1/\lambda_\alpha$.

It is observed that $\bar{G}^{ik}$ is the Eulerian component at x^i of the contravariant metric tensor at x_0^i , i.e.

$$\bar{G}^{ik}(x, t) = g^{mp}(x_0) \frac{\partial x^i}{\partial x_0^m} \frac{\partial x^k}{\partial x_0^p}; \quad (21)$$

the bar notation is necessary because the operation of raising a suffix does not commute with that of taking Eulerian components. It is convenient to regard the tensor $\bar{G}^{ik}$ as the fundamental measure of the deformation at x^i at time t when a Lagrangian point of view is adopted (i.e. when x^i is regarded as an explicit function of the x_0^i 's and t), just as G_{ik} is to be regarded as the natural measure of the deformation when an Eulerian point of view is adopted (i.e. when x_0^i is regarded as an explicit function of the x^i 's and t). These tensors will be referred to as the (contravariant and covariant) *deformation tensors*. The covariant and contravariant strain tensors defined by

$$e_{ik}(x, t) \equiv \frac{1}{2}[g_{ik}(x) - G_{ik}(x, t)], \quad \bar{e}^{ik}(x, t) \equiv \frac{1}{2}[\bar{G}^{ik}(x, t) - g^{ik}(x)] \quad (22)$$

may be useful concepts for particular problems, and may occur frequently in expressions for the free energy, but we can regard them as derived quantities which can always be expressed in terms of the deformation tensors. For the purpose of pursuing the general theory, greater simplicity can be achieved by using the deformation tensors (cf. §5).

As an example, we observe that the free-energy function for an isotropic elastic solid must always be reducible to the form

$$F = F(J_1, J_2, J_3, T, T_0, K_N), \quad (23)$$

$$\text{or, alternatively,} \quad F = F(I_1, I_2, I_3, T, T_0, K_N), \quad (24)$$

where the K_N 's are absolute scalar physical constants. The equations of state can then be written down by making use of the formulae (19) and (20); in the general case they are given by

$$p_{ik} = -2\rho \left\{ \frac{\partial F}{\partial J_1} G_{ik} + \frac{\partial F}{\partial J_2} (J_1 G_{ik} - G_{im} G_{km}^m) + \frac{\partial F}{\partial J_3} J_3 g_{ik} \right\}, \quad (25)$$

$$\text{or} \quad p^{ik} = 2\rho \left\{ \frac{\partial F}{\partial I_1} \bar{G}^{ik} + \frac{\partial F}{\partial I_2} (I_1 \bar{G}^{ik} - \bar{G}^{im} \bar{G}_m^k) + \frac{\partial F}{\partial I_3} I_3 g^{ik} \right\}, \quad (26)$$

$$\text{where} \quad \rho = \rho_0 J_3^{\frac{1}{3}} = \rho_0 I_3^{-\frac{1}{3}}. \quad (27)$$

The possible forms of F for an isotropic material are restricted by the requirement that $\partial F/\partial J_1 + 2\partial F/\partial J_2 + \partial F/\partial J_3$ (or equally $\partial F/\partial I_1 + 2\partial F/\partial I_2 + \partial F/\partial I_3$) must vanish when $J_1 = 3, J_2 = 3, J_3 = 1$ (or $I_1 = 3, I_2 = 3, I_3 = 1$) and $T = T_0$.

4. RELATION TO EARLIER WORK

All the quantities involved in the equations of state, in the present method of formulation, are tensors at a point of the deformed material. So far as the measures of deformation are concerned, there is a complete duality between the deformed and the undeformed configurations, and suitable definitions for deformation and strain

tensors at x_0^i can be written down simply by interchanging x^i and x_0^i , t and t_0 in the definitions given above. For example,

$$\Gamma_{ik}(x_0, t_0, t) \equiv g_{mp}(x) \frac{\partial x^m}{\partial x_0^i} \frac{\partial x^p}{\partial x_0^k} \quad (28)$$

is the analogue of $G_{ik}(x, t, t_0)$ and

$$\epsilon_{ik}(x_0, t_0, t) \equiv \frac{1}{2}[\Gamma_{ik}(x_0, t_0, t) - g_{ik}(x_0)], \quad (29)$$

the analogue of $-e_{ik}(x, t, t_0)$, is the tensor which is usually regarded as the strain tensor when a Lagrangian point of view is adopted. On the other hand, the stresses involved in the equations of state are those at x^i at the current time t , and they primarily form a tensor at x^i . It is therefore simplest to work with tensors at x^i at time t throughout.

The only comparable discussion of anisotropic elastic solids by previous writers appears to be that of Murnaghan (1937, §6), who applies the principle of virtual work to obtain equations relating the stress tensor at a point of the *deformed* material and measures of the deformation which are components of tensors at the *initial* position of the material considered. These equations involve the displacement functions explicitly.

Murnaghan gives a more detailed treatment of an isotropic elastic solid (1937, §3), in which all the tensors involved are tensors at a point of the deformed material, and in which the Eulerian point of view is adopted. In order to obtain Murnaghan's form of the equations of state, we suppose that F is expressed in the form

$$F = F(e_i^k, T, T_0, K_N), \quad (30)$$

which is seen to be possible for an isotropic material from equations (23) and (15) and the relation

$$e_i^k = \frac{1}{2}(\delta_i^k - G_i^k). \quad (31)$$

A direct application of (13), making use of (31) and (17), gives the equations of state in the form

$$p_{ik} = \frac{1}{2}\rho \left\{ (g_{im} - 2e_{im}) \frac{\partial F}{\partial e_m^k} + (g_{km} - 2e_{km}) \frac{\partial F}{\partial e_m^i} \right\}; \quad (32)$$

this reduces to the apparently unsymmetrical expression

$$p_k^i = \rho \left\{ \frac{\partial F}{\partial e_i^k} - 2e_m^i \frac{\partial F}{\partial e_m^k} \right\} \quad (33)$$

given by Murnaghan,* when the e_i^k 's occur only through the invariants J_1, J_2, J_3 . (It is noted that e_i^k and $e_{\cdot i}^k$ are not distinguished in the above; if both occur in a given expression for F , they are to be regarded as the same symbol (e_{ii}^k) in performing the differentiations.)

An independent derivation of the stress-strain relations for an isotropic elastic solid which is deformed isothermally is given by Rivlin (1948*b*), who adopts a

* Strictly, Murnaghan's derivatives are partial derivatives with respect to e_i^k holding the metric tensor components constant. It is possible to identify (33) and Murnaghan's equation because neither deformation nor metric tensor components occur in F , except in the combinations e_i^k .

Lagrangian point of view; the equations are intended for use with an orthogonal Cartesian co-ordinate system, and a tensor notation is not used. The need for distinguishing between tensors at different points is not apparent in a Cartesian system of reference, and it is interesting to note that Rivlin's equations are essentially relationships connecting p^{ik} with the $\bar{e}^{ik}$'s, not with the more usual Lagrangian strain components ϵ_{ik} . These equations can be deduced from (26) by substituting for $\bar{G}^{ik}$ in terms of the metric tensor and strain tensors, thus:

$$p^{ik} = 2\rho \left\{ \frac{\partial F}{\partial I_1} (g^{ik} + 2\bar{e}^{ik}) - I_3 \frac{\partial F}{\partial I_2} (g^{ik} - 2e^{ik}) + \left(I_2 \frac{\partial F}{\partial I_2} + I_3 \frac{\partial F}{\partial I_3} \right) g^{ik} \right\}. \quad (34)$$

If allowance is made for the differences of notation,* equation (34) is recognized as the invariant form of the stress-strain relations derived by Rivlin (1948*b*, equations (4.5)).

It has now been verified, incidentally, that Murnaghan's and Rivlin's equations of state for isotropic elastic solids are effectively the same.

The discussion by previous writers, concerning the finite straining of a continuum deformed elastically from an unstressed equilibrium position does not appear to include a derivation of compatibility conditions for the strains at different points of the material. The notation of the present paper lends itself very readily to the derivation of a necessary and sufficient compatibility condition.

5. A COMPATIBILITY CONDITION FOR FINITE STRAINS

An elastic continuum is now considered, in a particular deformed configuration. Any such configuration can be considered as attained at some instant t during actual motion, but the history of the material previous to the instant t is now supposed unknown, and the parameters t_0 and t are to be regarded as no more than convenient labels denoting the undeformed state and the given deformed state. We suppose that the deformation of the material at x^i is specified by a symmetric deformation tensor $G_{ik}(x, t)$, and we require a necessary and sufficient condition that the tensor field $G_{ik}(x, t)$ is consistent with the existence of displacement functions $x_0^i(x, t)$.

In order to find a necessary condition, we suppose that $G_{ik}(x, t)$ is the component of a physically significant deformation tensor, related to displacement functions $x_0^i(x, t)$ and the known metric tensor field by equation (4). If an arbitrary convected system of co-ordinates ξ^j is set up in the material, the metric tensor at x_0^i and the deformation tensor at x^i have the same convected-components $\gamma_{jl}(\xi, t_0)$ at a certain point of the material ξ^j . The Riemann-Christoffel tensor $\rho_{jlnq}(\xi, t_0)$, defined in terms of $\gamma_{jl}(\xi, t_0)$ as metric tensor and the co-ordinates ξ^j , must clearly vanish, i.e.

$$\rho_{jlnq}(\xi, t_0) = 0 \quad (\text{all } \xi^j). \quad (35)$$

If a special choice of convected co-ordinate system is now made, namely that for which the surfaces $\xi^j = \text{constant}$ coincide with the surfaces $x^i = \text{constant}$ at the instant t , we may write

$$\xi^j = x^j, \quad \gamma_{jl}(\xi, t_0) = G_{jl}(x, t), \quad \rho_{jlnq}(\xi, t_0) = R_{jlnq}(x, t), \quad (36)$$

* It is readily verified that Rivlin's ϵ_{xx} , etc., are the orthogonal Cartesian components of ϵ_{ik} ($i = k$) or $2\epsilon_{ik}$ ($i \neq k$); ϵ'_{xx} , etc., are similarly related to the tensor e^{ik} ; $\epsilon_{\xi\xi}$, etc., are similarly related to the tensor e_{ik} ; W is to be identified with $[\rho_0 F]_{T-T_0}$.

using the same suffixes in the fixed and convected co-ordinate systems. Equation (35) then becomes

$$R_{ikmp}(x, t) = 0 \quad (\text{all } x^i), \quad (37)$$

where

$$R_{ikmp} \equiv \frac{1}{2} \left(\frac{\partial^2 G_{km}}{\partial x^i \partial x^p} + \frac{\partial^2 G_{ip}}{\partial x^k \partial x^m} - \frac{\partial^2 G_{kp}}{\partial x^i \partial x^m} - \frac{\partial^2 G_{im}}{\partial x^k \partial x^p} \right) + \bar{G}^{jl} ([_{ip, j}] [_{km, l}] - [_{im, j}] [_{kp, l}]), \quad (38)$$

$$[_{ik, m}] \equiv \frac{1}{2} \left(\frac{\partial G_{km}}{\partial x^i} + \frac{\partial G_{im}}{\partial x^k} - \frac{\partial G_{ik}}{\partial x^m} \right), \quad (39)$$

and $\bar{G}^{ik}$ is the symmetric tensor defined by (14). The condition (37) is a necessary compatibility condition for the deformations at different points.

It can be shown to be also a sufficient condition, in the following way. We suppose that equation (37) is satisfied by a given symmetric tensor field $G_{ik}(x, t)$ associated with the material in the deformed state (t). This equation is recognized immediately as a sufficient condition that the functions $G_{jl}(y, t)$ of variables y^j are possible metric tensor components associated with co-ordinates y^j in Euclidean space. We can therefore set up a new fixed system of co-ordinates y^j , with the metric tensor $G_{jl}(y, t)$, which is arbitrary to the extent of a rigid-body translation and rotation in space. The point y^j in the new system of reference will have co-ordinates

$$x_0^i = x_0^i(y, t) \quad (40)$$

in the old. The metric tensor components in the new and old systems are related by the usual tensor transformation law

$$G_{jl}(y, t) = g_{mp}(x_0) \frac{\partial x_0^m}{\partial y^j} \frac{\partial x_0^p}{\partial y^l}. \quad (41)$$

Hence, making a simple change of notation from y^j to x^i , we see that there exist three functions $x_0^i(x, t)$ of the variables x^i (defined by (40)), which are related to the given functions $G_{ik}(x, t)$ and the known metric tensor field $g_{ik}(x)$ by equation (4). The tensor field $G_{ik}(x, t)$ is therefore physically significant as a deformation tensor field, for a particular deformation in which the material at x^i in the deformed state was at $x_0^i(x, t)$ in the undeformed state. The deformation consistent with the given deformation tensor field is arbitrary to the extent of a rigid-body translation and rotation.

From the symmetry properties of the Riemann-Christoffel tensor, there are six linearly independent equations (37), but of these only three are completely independent because of the three differential relations of the form

$$R_{ikpr, m} + R_{ikrm, p} + R_{ikmp, r} = 0 \quad (42)^*$$

(Bianchi's identity). Therefore the functions $G_{ik}(x, t)$ have three degrees of freedom, just as the displacement functions have three degrees.

* In this equation $R_{ikmp, r}$ denotes the Eulerian component at x^i of $\rho_{jlnq, s}(\xi, t_0)$, the co-variant derivative of $\rho_{jlnq}(\xi, t_0)$ at ξ^j at time t_0 , not the covariant derivative of $R_{ikmp}(x, t)$ at x^i .

The definition (38) could, if necessary, be expressed in terms of the strain tensor $e_{ik}(x, t)$, but with some loss of simplicity in form.* It is easily verified that, when the co-ordinate system is Cartesian, and when squares and products of the space derivatives of the e_{ik} 's can be neglected, equation (37) reduces to the familiar compatibility condition for infinitesimal strains

$$\frac{\partial^2 e_{km}}{\partial x^i \partial x^m} + \frac{\partial^2 e_{ip}}{\partial x^k \partial x^m} - \frac{\partial^2 e_{kp}}{\partial x^i \partial x^m} - \frac{\partial^2 e_{im}}{\partial x^k \partial x^p} = 0. \quad (43)$$

6. ALMOST INCOMPRESSIBLE MATERIALS

For the purpose of discussing materials in which the stresses required to change the volume of an element appreciably are much greater than those required to change its shape, it is convenient to define a *distortion tensor* G'_{ik} by

$$G'_{ik}(x, t) \equiv [J_3(x, t)]^{-\frac{1}{3}} G_{ik}(x, t). \quad (44)$$

The distortion tensor is a measure of the deformation of an element at x^i which would remain if all the initial linear dimensions of the element could be considered extended in the ratio $J_3^{-\frac{1}{3}}$ in order to make the initial and final densities equal. The corresponding quantities derived from G'_{ik} and G_{ik} are denoted by the same symbols with and without primes; for example, J'_3 is the determinant of the G'^k_i 's, which is identically unity, and $\bar{G}'^{ik}$ denotes the symmetric tensor reciprocal to G'_{ik} . Since J_3 and G'_{ik} are independent measures of the changes in size and shape, respectively, of a material element, we may regard F as an absolute scalar combination of the tensors

$$g_{ik}, G'_{ik}, J_3, T, T_0, K_N::i::k::\dots \quad (N = 1, 2, 3, \dots). \quad (45)$$

It follows from (13), (17) and (44) that the stresses are then given by

$$p^{ik} = \rho \left\{ \frac{\partial F}{\partial g_{ik}} + \frac{\partial F}{\partial g_{ki}} + \left(\frac{2}{3} G'_{mp} \frac{\partial F}{\partial G'_{mp}} - 2 J_3 \frac{\partial F}{\partial J_3} \right) g^{ik} \right\}. \quad (46)$$

It is usual to introduce, as an idealization of materials of small compressibility, the concept of an incompressible material (cf. Rivlin 1948*a*). Something in the nature of an external constraint is envisaged, which imposes the restriction $J_3 = 1$ on the possible deformations of any material element, so that the free energy may be regarded as a function F' of the tensors (45) with J_3 omitted. A material in which, at a given temperature, the stress at a point has six degrees of freedom while the deformation has only five is not elastic according to the definition of § 1. It is necessary to postulate that, in an incompressible material which behaves ideally elastically in distortion, the deformation and temperature determine the stress completely, *except for an arbitrary added isotropic pressure*. Then the equations of state (46) are modified to read

$$p^{ik} = p'^{ik} - p'' g^{ik}, \quad (47)$$

* It is equally possible to express the compatibility condition as the vanishing of a tensor $P_{ikmp}(x_0, t)$ at x_0^i , whose components are defined by equations of the type (38) and (39) with x^i replaced by x_0^i and G_{ik} by Γ_{ik} . The arguments of § 5 remain valid if the deformed and undeformed states are interchanged.

where p'' is an arbitrary function of position and time, and

$$p'^{ik} = \rho \Delta'^{ik} F', \quad (48)$$

the operator Δ'^{ik} being defined in the same way as Δ^{ik} except that the G'_{ik} 's are held constant, instead of the G_{ik} 's, in performing the differentiations.

The use of equations (47) and (48) as an approximate form of the equations of state of an ideally elastic material of small compressibility requires justification. It is of interest to examine what assumptions must be made regarding the form of the free-energy function in order that equations of this type should give a reliable approximation to the equations of state. The property of being almost incompressible must be regarded as a consequence of the form of the function F ; essentially, F must be very sensitive to changes in J_3 , so that any appreciable (positive or negative) dilatation, at constant distortion and temperature, gives rise to large restoring forces. Approximate equations of state must be based on this type of free-energy function, not on one which is independent of J_3 as in the above derivation of (47). The following argument, although not rigorous, indicates the kind of assumptions which must be introduced.

For given values of T and G'_{ik} , F is a function of J_3 , with a minimum value F' at, say, $J_3 = 1 + \epsilon$. In a material of small compressibility (at constant temperature and constant distortion), $|J_3 - 1 - \epsilon|$ will be small at all finite stresses, and we suppose that F may be expanded as a power series, thus

$$F = F' + (J_3 - 1 - \epsilon)^2 F'' + (J_3 - 1 - \epsilon)^3 F''' + (J_3 - 1 - \epsilon)^4 F'''' + \dots \quad (49)$$

In general ϵ , F' , F'' , F''' , etc., will be scalar combinations of g_{ik} , G'_{ik} , T , T_0 and physical constants; ϵ is not necessarily negligibly small. Introducing the abbreviation

$$\Delta''^{ik} \equiv \left\{ \frac{\partial}{\partial g_{ik}} + \frac{\partial}{\partial g_{ki}} + \frac{2}{3} g^{ik} G'_{mp} \frac{\partial}{\partial G'_{mp}} \right\}, \quad (50)$$

the equations of state (46) may be written as

$$p^{ik}/\rho = \Delta''^{ik} F' + \{(J_3 - 1 - \epsilon)^2 \Delta''^{ik} F'' + (J_3 - 1 - \epsilon)^3 \Delta''^{ik} F''' + \dots\} \\ - \{2(J_3 - 1 - \epsilon) F'' + 3(J_3 - 1 - \epsilon)^2 F''' + \dots\} (\Delta''^{ik} \epsilon + 2J_3 g^{ik}). \quad (51)$$

Each term of the first series in braces in (51) is of a smaller order of magnitude than the corresponding term in the second. If the function F'' (a measure of 'bulk modulus') takes sufficiently large values at all attainable distortions to make $|J_3 - 1 - \epsilon|$ negligible compared with unity for all attainable stresses, the first series in braces may be neglected. It is then only necessary to introduce the assumption that $g_{im} \Delta''^{mk} \epsilon$ is negligibly small compared with unity, in order to write (51) in the form of (47) and (48), in which F' is a function of the state which is independent of J_3 . But, unless further assumptions are introduced, it is necessary to calculate p'' from the formula

$$p'' = 2\rho \left\{ J_3 \frac{\partial F'}{\partial J_3} - \frac{1}{3} G'_{mp} \frac{\partial F'}{\partial G'_{mp}} \right\}. \quad (52)$$

If the material is such that ϵ is always negligibly small, so that no appreciable change in volume can be produced by any system of stresses, and if the boundary

conditions in a particular problem do not demand an *exact* evaluation of volume changes, it will be possible to write $J_3 = 1$ as a sufficient approximation for all purposes except the calculation of p'' from (52). Under such conditions it will in fact be unnecessary to calculate p'' in this way, since the small variations in J_3 which must occur from place to place in the deformed material will be just sufficient to make the function $p''(x, t)$ consistent with the equations of motion. (The equation of continuity, three equations of motion, six equations (48) and the condition $J_3 = 1$ may be regarded as eleven equations for the eleven unknown functions of position and time, ρ , p'' , six components p'^{ik} and three displacement functions, as in the case of an ideally incompressible material.) Under these conditions, the function F' is obtained, to a sufficient degree of approximation, by substituting $J_3 = 1$ in F .

The generalized form of (48), applicable when F is expressed in terms of J_3 (or I_3), a set of tensors $h'_{M:i:k} \dots (M = 1, 2, 3, \dots)$ measuring the distortion (each of which is to be regarded as a combination of the tensors g_{ik} , G'_{ik} and universal constants), the temperature and physical constant tensors, is

$$p'^{ik} = \rho \sum_M \frac{\partial F'}{\partial h'_{M:i:k} \dots} \Delta'^{ik} h'_{M:i:k} \dots \quad (53)$$

In particular, in the case of an isotropic material with a free-energy function expressed in terms of the invariants* I'_1 , J'_1 and J_3 , we have

$$p'^{ik} = 2\rho \left\{ \frac{\partial F'}{\partial I'_1} \bar{G}'^{ik} - \frac{\partial F'}{\partial J'_1} G'^{ik} \right\}. \quad (54)$$

As an example, the expression proposed by Mooney (1940) for the strain-energy function of an incompressible rubber subjected to isothermal deformations is equivalent to

$$F' = (C_1 I'_1 + C_2 J'_1) / \rho_0 + \text{a constant}, \quad (55)$$

where C_1 and C_2 are absolute scalar physical constants. The corresponding invariant form of the stress-deformation relations is

$$p^{ik} = 2(\rho/\rho_0) (C_1 \bar{G}'^{ik} - C_2 G'^{ik}) - p'' g^{ik}, \quad (56)$$

where p'' is an arbitrary function of position and time to be determined by the equations of motion. On substituting $J_3 = 1$, the (approximate) equations of state for rubber become

$$p^{ik} = 2C_1 \bar{G}^{ik} - 2C_2 G^{ik} - p'' g^{ik}. \quad (57)$$

7. THE TORSION OF AN ANISOTROPIC CYLINDER

A simple illustration will now be given of the derivation and use of equations of state involving anisotropic constants. For simplicity, a problem is selected in which the boundary conditions do not involve the stresses explicitly. We shall calculate the forces at the boundaries which are required to maintain a right circular cylinder of a particular type of (almost) incompressible anisotropic elastic solid in a state of

* In terms of the principal extension ratios, we have $I'_1 = J'_2 = \sum_{\alpha} \lambda_{\alpha}^2$, $I'_2 = J'_1 = \sum_{\alpha} \lambda_{\alpha}^{-2}$ in the case of an incompressible material.

simple torsion in the absence of body forces. This problem has been discussed quite generally by Rivlin (1948*b*), in the case when isotropy can be assumed.

We suppose that the free energy per unit mass is given by

$$F = B_i^k \bar{G}_k^i + C_i^k G_k^i + f(J_3), \quad (58)$$

where the anisotropic physical constants B_i^k and C_i^k are mixed, second-order absolute tensors, arbitrary except that B_{ik} and C_{ik} may be assumed symmetric without loss of generality. In the undeformed state all elements of the cylinder will be assumed similarly oriented, so far as their directional properties are concerned. We shall find that there are restrictions on the possible nature of the constants, arising (i) from the requirement that the stresses vanish in the undeformed state and (ii) from the implicit assumption that an equilibrium state of simple torsion can be maintained by means of forces at the boundaries. Both types of restriction permit the isotropic form of (58), obtained by writing $B_i^k = \frac{1}{3} B_m^m \delta_i^k$, $C_i^k = \frac{1}{3} C_m^m \delta_i^k$, which corresponds to a function F' of the form (55). The anisotropic material represented by (58) may therefore be regarded as a generalization of Mooney's prototype for rubber.

From equations (13), (17), (18), (19) and (44), the exact stress-deformation relations corresponding to (58) are

$$p^{ik} = \rho B_m^r (\delta_r^i \bar{G}'^{km} + \delta_r^k \bar{G}'^{im} - \frac{2}{3} \bar{G}_r'^m g^{ik}) - \rho C_m^r (g^{im} G_r'^k + g^{km} G_r'^i - \frac{2}{3} G_r'^m g^{ik}) - 2\rho J_3 f'(J_3) g^{ik}. \quad (59)$$

The stresses vanish automatically when $J_3 = 1$ and $G_{ik}' = g_{ik}$, only if

$$f'(1) = 0, \quad B_i^k - C_i^k = \frac{1}{3} (B_m^m - C_m^m) \delta_i^k. \quad (60)$$

Since the material is almost incompressible and the associated function ϵ clearly vanishes, we may assume $J_3 = 1$ except for the purpose of calculating $f'(J_3)$ (which will take finite values of the same sign as $J_3 - 1$ when J_3 differs from unity by a very small amount). Then the equations of state may be written in the form (47), where

$$p'^{ik} = \rho_0 (B_m^i \bar{G}^{km} + B_m^k \bar{G}^{im}) - \rho_0 g^{ir} g^{kl} (C_r^m G_{lm} + C_l^m G_{rm}) \quad (61)$$

(approximately) and p'' is to be determined as a function of position from the equations of equilibrium.

Referred to a system of cylindrical polar co-ordinates r, ϕ, z , such that the solid cylinder is bounded by the surface $r = a$, a simple torsion is defined as a deformation in which (to a degree of approximation corresponding to the neglect of changes in volume) the material at (r, ϕ, z) in the deformed state was at $(r, \phi - \psi z, z)$ in the undeformed state, where ψ is a constant. It is more convenient to work in terms of orthogonal Cartesian co-ordinates $x^i = x, y, z (= r \cos \phi, r \sin \phi, z)$, so that the relations between the initial and final co-ordinates of the same material are

$$x_0 = x \cos \psi z + y \sin \psi z, \quad y_0 = -x \sin \psi z + y \cos \psi z, \quad z_0 = z, \quad (62)$$

$$\text{or} \quad x = x_0 \cos \psi z_0 - y_0 \sin \psi z_0, \quad y = x_0 \sin \psi z_0 + y_0 \cos \psi z_0, \quad z = z_0. \quad (63)$$

The Cartesian components b_i^k, c_i^k of the physical constants associated with the material at x_0^i in the undeformed state are supposed known numerical constants.

independent of x_0^i ; B_i^k , C_i^k are the Eulerian components at x^i of the tensors b_i^k , c_i^k at x_0^i . It is therefore necessary to make the following substitutions in (61), in order to obtain the Cartesian stress components:

$$g^{ik} = \delta_i^k, \quad \bar{G}^{ik} = \frac{\partial x^i}{\partial x_0^m} \frac{\partial x^k}{\partial x_0^m}, \quad G_{ik} = \frac{\partial x_0^m}{\partial x^i} \frac{\partial x_0^m}{\partial x^k}, \quad B_i^k = \frac{\partial x_0^m}{\partial x^i} \frac{\partial x^k}{\partial x_0^p} b_m^p, \quad C_i^k = \frac{\partial x_0^m}{\partial x^i} \frac{\partial x^k}{\partial x_0^p} c_m^p, \quad (64)$$

where

$$\frac{\partial x_0^i}{\partial x^k} = \begin{pmatrix} \cos \psi z & \sin \psi z & \psi(-x \sin \psi z + y \cos \psi z) \\ -\sin \psi z & \cos \psi z & \psi(-x \cos \psi z - y \sin \psi z) \\ 0 & 0 & 1 \end{pmatrix}, \quad \frac{\partial x^i}{\partial x_0^k} = \begin{pmatrix} \cos \psi z & -\sin \psi z & -\psi y \\ \sin \psi z & \cos \psi z & \psi x \\ 0 & 0 & 1 \end{pmatrix}. \quad (65)^*$$

The full expressions for the stresses are omitted in the interests of conciseness. On substituting them in the equations of equilibrium, in the form

$$\partial p'' / \partial x^i = \partial p_i'^k / \partial x^k, \quad (66)$$

it is found that there is no solution for $p''(x, y, z)$ unless

$$2b_2^3 + c_2^3 = 0, \quad 2b_3^1 + c_3^1 = 0. \quad (67)$$

Taking into account also the restriction (60) on the physical constants, in the transformed form

$$b_i^k - c_i^k = \frac{1}{3}(b_m^m - c_m^m) \delta_i^k \equiv b \delta_i^k, \quad (68)$$

it is seen that an equilibrium state of simple torsion is impossible unless

$$b_i^k = \begin{pmatrix} c_1^1 + b & c_1^2 & 0 \\ c_1^2 & c_2^2 + b & 0 \\ 0 & 0 & c_3^3 + b \end{pmatrix}, \quad c_i^k = \begin{pmatrix} c_1^1 & c_1^2 & 0 \\ c_1^2 & c_2^2 & 0 \\ 0 & 0 & c_3^3 \end{pmatrix}. \quad (69)$$

For a material characterized by physical constants with these Cartesian components at x_0^i , a state of simple torsion is possible, provided that the stresses in the deformed material have Cartesian components p_{ik} given by

$$\begin{aligned} \frac{p_{ik}}{\rho_0} = & c_1^1 \begin{pmatrix} 0 & 0 & \psi(x \sin \psi z - y \cos \psi z) \cos \psi z \\ \dots & 0 & \psi(x \sin \psi z - y \cos \psi z) \sin \psi z \\ \dots & \dots & -\psi^2(x \sin \psi z - y \cos \psi z)^2 \end{pmatrix} + c_1^2 \begin{pmatrix} 0 & 0 & \psi(x \cos 2\psi z + y \sin 2\psi z) \\ \dots & 0 & \psi(x \sin 2\psi z - y \cos 2\psi z) \\ \dots & \dots & -\psi^2[(x^2 - y^2) \sin 2\psi z - 2xy \cos 2\psi z] \end{pmatrix} \\ & + c_2^2 \begin{pmatrix} 0 & 0 & -\psi(x \cos \psi z + y \sin \psi z) \sin \psi z \\ \dots & 0 & \psi(x \cos \psi z + y \sin \psi z) \cos \psi z \\ \dots & \dots & -\psi^2(x \cos \psi z + y \sin \psi z)^2 \end{pmatrix} + (c_3^3 + b) \begin{pmatrix} \psi^2 y^2 & -\psi^2 xy & -\psi y \\ \dots & \psi^2 x^2 & \psi x \\ \dots & \dots & 0 \end{pmatrix} \\ & + \left(b - \frac{p''}{2\rho_0}\right) \begin{pmatrix} 1 & 0 & 0 \\ \dots & 1 & 0 \\ \dots & \dots & 1 \end{pmatrix}, \quad (70) \end{aligned}$$

where p'' is given by

$$\begin{aligned} -p''/(\rho_0 \psi^2) = & (c_2^2 - c_1^1) [(x^2 - y^2) \cos 2\psi z + 2xy \sin 2\psi z] \\ & + 2c_1^2 [(x^2 - y^2) \sin 2\psi z - 2xy \cos 2\psi z] + (c_3^3 + b)(x^2 + y^2) + \text{a constant}. \quad (71) \end{aligned}$$

* With the convention that i refers to the row, k to the column of the array.

If we now consider the Cartesian axes rotated through an angle ϕ , so that their directions coincide with the polar co-ordinate directions at an arbitrarily chosen point $(r \cos \phi, r \sin \phi, z)$, we obtain the following expressions for the physical components of stress in the polar co-ordinate directions:

$$\left. \begin{aligned} \widehat{rr} &= \rho_0 \psi^2 r^2 [(c_2^2 - c_1^1) \cos 2(\phi - \psi z) - 2c_1^2 \sin 2(\phi - \psi z) + c_3^3 + b] - P, \\ \widehat{\phi\phi} &= \rho_0 \psi^2 r^2 [(c_2^2 - c_1^1) \cos 2(\phi - \psi z) - 2c_1^2 \sin 2(\phi - \psi z) + 3(c_3^3 + b)] - P, \\ \widehat{zz} &= \rho_0 \psi^2 r^2 [-c_1^1 - c_2^2 + c_3^3 + b] - P, \\ \widehat{\phi z} &= 2\rho_0 \psi r [c_1^1 \sin^2 (\phi - \psi z) + c_2^2 \cos^2 (\phi - \psi z) - c_1^2 \sin 2(\phi - \psi z) + c_3^3 + b], \\ \widehat{zr} &= \rho_0 \psi r [(c_2^2 - c_1^1) \sin 2(\phi - \psi z) + 2c_1^2 \cos 2(\phi - \psi z)], \\ \widehat{r\phi} &= 0. \end{aligned} \right\} \quad (72)$$

The stresses are determined except for a superposed arbitrary constant isotropic pressure P .

It is concluded that, so long as the axis of the cylinder is a principal direction for the physical constant tensors at a point of the undeformed material, an equilibrium state of simple torsion can be maintained by means of suitable stress distributions applied to the plane ends and to the curved surface of the cylinder. The surface stresses over the ends $z = \text{constant}$ are given by the expressions (72) for $\widehat{zr}$, $\widehat{\phi z}$, $\widehat{zz}$, and the stresses required on the curved surface $r = a$ are given by the expressions for $\widehat{rr}$, $\widehat{r\phi}$, $\widehat{zr}$. In the case of an isotropic solid the stresses on the curved surface can be made to vanish by choice of P . The expressions for the stresses on the plane ends are then seen to agree with those obtained by Rivlin (1948*b*, equations (14, 16)) for isotropic materials of this type, if we write

$$B_i^k = b_i^k = C_1 \delta_i^k / \rho_0, \quad C_i^k = c_i^k = C_2 \delta_i^k / \rho_0, \quad b = (C_1 - C_2) / \rho_0, \quad P = C_1 \psi^2 a^2,$$

in order to make the notations agree.

REFERENCES

- Mooney, M. 1940 *J. Appl. Phys.* **11**, 582.
Murnaghan, F. D. 1937 *Amer. J. Math.* **59**, 235.
Oldroyd, J. G. 1950 *Proc. Roy. Soc. A*, **200**, 523.
Rivlin, R. S. 1948*a* *Phil. Trans. A*, **240**, 459.
Rivlin, R. S. 1948*b* *Phil. Trans. A*, **241**, 379.

Contributions to the theory of heat transfer through a laminar boundary layer

BY M. J. LIGHTHILL

Department of Mathematics, University of Manchester

(Communicated by S. Goldstein, F.R.S.—Received 9 February 1950)

An approximation to the heat transfer rate across a laminar incompressible boundary layer, for arbitrary distribution of main stream velocity and of wall temperature, is obtained by using the energy equation in von Mises's form, and approximating the coefficients in a manner which is most closely correct near the surface. The heat transfer rate to a portion of surface of length l (measured downstream from the start of the boundary layer) and unit breadth is given as

$$-\frac{\frac{1}{2}k}{(\frac{1}{2})!} \left(\frac{3\sigma\rho}{\mu^2} \right)^{\frac{1}{2}} \int_0^l \left(\int_x^l \sqrt{\tau(z)} dz \right)^{\frac{1}{2}} dT_0(x),$$

where k is the thermal conductivity of the fluid, σ its Prandtl number, ρ its density, μ its viscosity, $\tau(x)$ is the skin friction, and $T_0(x)$ the excess of wall temperature over main stream temperature. A critical appraisalment of the formula (§3) indicates that it should be very accurate for large σ , but that for σ of order 0.7 (i.e. for most gases) the constant $\frac{1}{2}3^{\frac{1}{2}}/(\frac{1}{2})! = 0.807$ should be replaced by 0.73, when the error should not exceed 8 % for the laminar layers that occur in practical aerodynamics. This yields a formula $Nu = 0.52\sigma^{\frac{1}{2}}(R\sqrt{C_f})^{\frac{1}{2}}$ for Nusselt number in terms of the Reynolds number R and the mean square root of the skin friction coefficient C_f , in the case of uniform wall temperature.

However, for the boundary layer with uniform main stream, the original formula is accurate to within 3% even for $\sigma = 0.7$. By known transformations an expression is deduced for heat transfer to a surface, with arbitrary temperature distribution along it, and with a uniform stream outside it at arbitrary Mach number (equation (42)). From this, the temperature distribution along such a surface is deduced (§4) in the case (of importance at high Mach numbers) when heat transfer to it is balanced entirely by radiation from it. This calculation, which includes the solution of a non-linear integral equation, gives higher temperatures near the nose, and lower ones farther back (figure 2), than are found from a theory which assumes the wall temperature uniform and averages the heat transfer balance. This effect will be considerably mitigated for bodies of high thermal conductivity; the author is not in a position to say whether or not it will be appreciable for metal projectiles. But for stony meteorites at a certain stage of their flight through the atmosphere it indicates that melting at the nose and re-solidification farther back may occur, for which the shape and constitution of a few of them affords evidence.

An appendix shows how the method for approximating and solving von Mises's equation could be used to determine the skin friction as well as heat transfer rate, but this line seems to have no advantage over established approximate methods.

1. INTRODUCTION

In the steady flight of an aircraft, projectile, or other body through the atmosphere, the front portions of the boundary layer on the surface are normally laminar.† Heat transfer across the boundary layer from a deliberately heated wall may be necessary, in a cold atmosphere, to prevent icing. Alternatively, heat transfer to the surface from a boundary layer heated by viscous stresses may, at high supersonic speeds, raise the body temperature to most undesirable levels. A brief account of the present state of the theory of heat transfer across a laminar boundary layer, first when the

† There is strong evidence that this remains true even at high supersonic speeds.

Mach number is small and then in the absence of that restriction, will be given before the contribution of this paper is described.

At low Mach numbers there is a range of exact solutions (Falkner & Skan 1930; Hartree 1937) of the boundary-layer equations for the velocity components—in which these have similar profiles at different stations along the surface—each obtainable by numerically solving an ordinary differential equation. The main-stream velocity is given by $u_1 = cx^m$, where x is measured along the surface from the start of the boundary layer; m may be any number exceeding -0.0904 , this being the condition for no boundary-layer separation (which occurs at once if at all). The velocity u parallel to the surface depends only on $\eta = (u_1/\nu x)^{\frac{1}{2}} y$, where y represents perpendicular distance from the wall and ν the kinematic viscosity. In fact it and the velocity v normal to the wall may be written in the form

$$u = u_1 f'(\eta), \quad v = -\frac{1}{2}(u_1 \nu/x)^{\frac{1}{2}} [(m+1)f(\eta) + (m-1)\eta f'(\eta)]. \quad (1)$$

The case $m = 0$ is the famous Blasius velocity distribution for the boundary layer at constant pressure. Now, neglecting terms of order the square of the Mach number,† the temperature T satisfies the equation

$$u \frac{\partial T}{\partial x} + v \frac{\partial T}{\partial y} = \frac{\nu}{\sigma} \frac{\partial^2 T}{\partial y^2}, \quad (2)$$

where $\nu/\sigma = \kappa$ is the diffusivity of heat ($= k/\rho c_p$ in terms of thermal conductivity k , density ρ and specific heat at constant pressure c_p), so that σ is the Prandtl number (0.71 for air). A solution of equation (2) taking the value T_1 (temperature of the main stream) for $y = \infty$, and taking a *constant* value T_0 along the wall $y = 0$, is easily obtained (when u, v are given by (1)) by assuming that T is a function of η only.

This is

$$\frac{T_1 - T}{T_1 - T_0} = \alpha(m, \sigma) \int_0^\eta \left[\exp \left(-\frac{(m+1)\sigma}{2} \int_0^\eta f d\eta \right) \right] d\eta, \quad (3)$$

where

$$\alpha(m, \sigma) = \left\{ \int_0^\infty \left[\exp \left(-\frac{(m+1)\sigma}{2} \int_0^\eta f d\eta \right) \right] d\eta \right\}^{-1} \quad (4)$$

has been calculated for a large number of pairs of values for m and σ . For $m = 0$, $\alpha = 0.332\sigma^{\frac{1}{2}}$ is a good interpolation formula (Pohlhausen 1921); for $m = 1$, $\alpha = 0.57\sigma^{0.4}$ is close for moderate σ but is an overestimate for $\sigma > 5$. Other known values are $\alpha(\frac{1}{9}, 0.7) = 0.331$, $\alpha(\frac{1}{3}, 0.7) = 0.384$, $\alpha(4, 0.7) = 0.813$. The rate of heat transfer to the wall per unit time per unit area at the point $(x, 0)$ is

$$Q(x) = \alpha(m, \sigma) k(T_1 - T_0) (u_1/\nu x)^{\frac{1}{2}}, \quad (5)$$

and one can display the heat transfer rate to an area of surface of length l and breadth b , in terms of a Nusselt number

$$\text{Nu} = \left(\int_0^l Q(x) b dx \right) / \left(\frac{k(T_1 - T_0)}{l} bl \right) \quad (6)$$

and a Reynolds number $R = u_1(l)/\nu = cl^{m+1}/\nu$, by the equation

$$\text{Nu } R^{-\frac{1}{2}} = \frac{\alpha(m, \sigma)}{\frac{1}{2}(m+1)}. \quad (7)$$

† And so neglecting, for example, the heat generated by friction.

Fage & Falkner (1931) made an approximate extension of this work to the case where the wall temperature T_0 is not a constant, but differs from its value at $x = 0$ by a multiple of x^n for some n . In equation (2) they replace u and v by their asymptotic forms as $y \rightarrow 0$, on the argument that the crucial part of the temperature profile will be near the wall, since the precise manner in which it asymptotes to its main stream value will not greatly affect the scale of curve required to attain the specified temperature at the wall. Series solutions in powers of y are then derived, and arbitrary coefficients therein are fixed by calculating the solutions for large y and comparing them with the main-stream values. The method, when applied to the case $n = 0$ (uniform wall temperature) and $\sigma = 0.77$, gave $\alpha(m, \sigma) = 0.31, 0.45$ and 0.60 for $m = 0, \frac{1}{3}$ and 1 respectively. The exact values are $0.303, 0.392$ and 0.513 .

For the Blasius boundary layer ($m = 0$) Chapman & Rubesin (1949) obtained exact solutions for T when the difference between wall temperature T_0 and stream temperature T_1 is given by a series $\sum_0^\infty a_n x^n$ or polynomial. Modifying their notation,† the solution takes the form $T - T_1 = \sum_0^\infty a_n x^n X_n(\eta)$, where the functions $X_n(\eta)$ are calculated with $\sigma = 0.72$ (each by numerically solving an ordinary differential equation) for $n = 0, 1, 2, 3, 4, 5$ and 10 . The local heat transfer rate to the wall $Q(x)$ is given by

$$Q(x) = k(u_1/\nu x)^{\frac{1}{2}} \sum_0^\infty a_n x^n X'_n(0), \quad (8)$$

where $-X'_n(0)$ is $0.296, 0.489, 0.597, 0.684, 0.744, 0.799$ and 1.006 respectively in the above seven cases.

The principal remaining calculations (Squire 1942) on laminar heat transfer at low Mach number, due to temperature difference between wall and stream, used an approximate method on the problem in which the main stream is arbitrary but the wall temperature is uniform. The velocity profile was taken as that appropriate to the Blasius solution, but the displacement thickness was allowed to vary with x in a manner appropriate to the assumed main stream. The temperature profile was taken similar to the velocity profile, but with scale in the y -direction altered in a manner to be determined from the energy integral equation.

The heating of the boundary layer due to friction was calculated for the Blasius velocity distribution, corresponding to uniform main-stream velocity, at low Mach numbers, by Pohlhausen (1921), who solved the relevant equation

$$u \frac{\partial T}{\partial x} + v \frac{\partial T}{\partial y} - \frac{\nu}{\sigma} \frac{\partial^2 T}{\partial y^2} = \frac{\nu}{c_p} \left(\frac{\partial u}{\partial y} \right)^2, \quad (9)$$

with $T = T_1$ in the main stream and with $\partial T/\partial y = 0$ at the wall $y = 0$. This determines the 'final state' in which the wall has been so heated or cooled by convection that no further heat transfer is taking place. The wall temperature T_0 takes the value $T_0 = T_1 [1 + \frac{1}{2}(\gamma - 1) M^2 \beta(\sigma)]$, where γ is the ratio of the specific heats, M the Mach number,‡ and $\beta(\sigma)$ was determined by numerical integration for various values of

† Their η is half of that used here and defined above.

‡ So that $(\gamma - 1) M^2 = u_1^2/c_p T_1$.

σ , and was found not to differ significantly from $\sigma^\dagger$ for moderate values. The same result, with $\beta = 1$, is of course true for arbitrary main-stream velocity distribution and Mach number for any gas for which $\sigma = 1$. The heating due to friction is seen to be small at low Mach numbers, but if desired it may be taken into account in solving a problem *with* heat transfer across the Blasius layer by observing that the *difference* in temperature distribution from that with zero heat transfer must satisfy equation (2), not equation (9). Thus, in problems which have been solved neglecting frictional heating, accuracy can be improved by replacing the difference in temperature between wall and stream by that between the wall and the calculated 'final' temperature.

When the Mach number M is not small, calculations (apart from those assuming $\sigma = 1$ and zero heat transfer) have been confined to the case of the laminar layer with uniform main-stream velocity. This was considered in great detail by Crocco (1946). He obtained certain solutions, fairly simply, by replacing the viscosity-temperature curve by a straight line $\mu/T = \text{constant}$. He compared these with solutions obtained with great labour using more exact curves, and found that the approximation involved in setting $\mu \propto T$ was a good one. Chapman & Rubesin (1949) make still simpler the method of attack when $\mu \propto T$ by using von Mises's idea of taking the stream function ψ as an independent variable.† For in variables (x, ψ) , where ψ represents the rate of mass flow between a given point and the wall, so that $\partial\psi/\partial x = -\rho v$, $\partial\psi/\partial y = \rho u$, the momentum and energy equations for the boundary layer at constant pressure, namely,‡

$$u \frac{\partial u}{\partial x} + v \frac{\partial u}{\partial y} = \frac{1}{\rho} \frac{\partial}{\partial y} \left(\mu \frac{\partial u}{\partial y} \right), \quad u \frac{\partial T}{\partial x} + v \frac{\partial T}{\partial y} = \frac{\mu}{\rho c_p} \left(\frac{\partial u}{\partial y} \right)^2 + \frac{1}{\rho \sigma} \frac{\partial}{\partial y} \left(\mu \frac{\partial T}{\partial y} \right), \quad (10)$$

become
$$\frac{\partial u}{\partial x} = \frac{\partial}{\partial \psi} \left(\mu \rho u \frac{\partial u}{\partial \psi} \right), \quad \frac{\partial T}{\partial x} = \frac{\mu \rho u}{c_p} \left(\frac{\partial u}{\partial \psi} \right)^2 + \frac{1}{\sigma} \frac{\partial}{\partial \psi} \left(\mu \rho u \frac{\partial T}{\partial \psi} \right). \quad (11)$$

In these equations the only terms which mark the influence of compressibility are μ and ρ , which both vary with T in the boundary layer although the pressure is uniform therein. However, their product is a constant if $\mu \propto T$, and it is only this product which occurs in (11). Hence, for given temperature distribution along the wall, u and T must be the same functions of x and ψ as they are at low Mach number.§ This does not mean that their profiles do not change with Mach number, for the meaning of ψ in terms of y is different. But the skin friction τ and heat transfer rate Q are unaltered. For, if zero suffix indicates values at the wall, then $u \sim (\tau/\mu_0) y$ and $\rho \rightarrow \rho_0$ as $y \rightarrow 0$, so $\psi = \int_0^y \rho u dy \sim \frac{1}{2}(\rho_0 \tau/\mu_0) y^2$. Hence

$$\tau = \frac{1}{2} \mu_0 \rho_0 \left[\left(\frac{\partial u}{\partial \psi^\dagger} \right)_0 \right]^2, \quad (12)$$

† They also emphasize the advantage of taking the constant μ/T to be, not its real value in the main stream, but a rough average of its real surface values.

‡ Here loss of energy of the gas by radiation is neglected. This is accepted in the work of §4 because radiation from the solid surface is taken into account. The former is difficult to estimate but is assumed to be trivial in comparison with the latter at all temperatures.

§ Except in so far as the constant $u_1^2/c_p T_1$ depends on Mach number, being $(\gamma - 1) M^2$.

and therefore is the same as at low Mach number, since $\mu\rho$ is a constant and u is the same function of x, ψ in the two flows. Similarly,

$$Q = k_0 \left(\frac{\partial T}{\partial y} \right)_0 = \frac{\mu_0}{c_p \sigma} \left(\frac{\partial T}{\partial \psi^{\frac{1}{2}}} \right)_0 \left(\frac{\partial \psi^{\frac{1}{2}}}{\partial y} \right)_0 = \frac{\mu_0 \rho_0}{2 c_p \sigma} \left(\frac{\partial T}{\partial \psi^{\frac{1}{2}}} \frac{\partial u}{\partial \psi^{\frac{1}{2}}} \right)_0 \quad (13)$$

is unchanged.

This result, that for a main stream with uniform speed but arbitrary Mach number, assuming $\mu \propto T$, the skin friction and heat transfer rate at any point are given by the low-speed equations, namely, by the Blasius velocity distribution and by equation (9) for the temperature, adds greatly to the importance of equation (9). Pohlhausen's solution

$$T = T_1 \text{ in main stream, } T = T_1 [1 + \frac{1}{2}(\gamma - 1) \sigma^{\frac{1}{2}} M^2] \text{ at wall, } Q = 0, \quad (14)$$

for the state with zero heat transfer is now seen to be valid for arbitrary M ; and other solutions may be obtained by adding to it an arbitrary solution of equation (2) with the Blasius values for the velocities u, v . For example, the solution with uniform wall temperature,

$$T = 0 \text{ in main stream, } T = -T_2 \text{ at wall, } Q = 0.332 \sigma^{\frac{1}{2}} k T_2 (u_1/\nu x)^{\frac{1}{2}}, \quad (15)$$

may be added to (14) to get the general solution for arbitrary Mach number and uniform wall temperature,

$$\left. \begin{aligned} T &= T_1 \text{ in main stream, } T = T_0 \text{ at wall,} \\ Q &= 0.332 \sigma^{\frac{1}{2}} k [T_1 (1 + \frac{1}{2}(\gamma - 1) \sigma^{\frac{1}{2}} M^2) - T_0] (u_1/\nu x)^{\frac{1}{2}} \end{aligned} \right\} \quad (16)$$

But solutions with non-uniform wall-temperature distributions can equally be added to (14).

The work of this paper is prompted by two considerations. First, the need to integrate the theory of heat transfer in laminar flow at low Mach number, by producing an approximate but simple formula for calculating it under arbitrary conditions of main-stream velocity and wall temperature. Secondly, the need to calculate the *true* final surface temperature distribution at high Mach number, not with zero heat transfer but with a positive value of the heat transfer rate at each point, corresponding approximately to that at which the surface loses heat by radiation. This can only be attempted for the boundary layer at constant main-stream velocity (such as obtains on a body with a wedge-shaped head). But, as §4 will show, the wall temperature in this final state is non-uniform, so that a theory applying under this condition is needed. Now the work of Chapman & Rubesin (1949) described above, though exact, is not suitable for this investigation, since it requires that the wall-temperature distribution be expressible as a polynomial or as a Maclaurin series; whereas the Maclaurin series for the true distribution which will be found in §4 both contains fractional powers of x and has a very small radius of convergence. In any case the formula of this paper is both especially simple and especially accurate for the constant-pressure layer, agreeing with the results of Chapman & Rubesin (1949) to within 3 %.

The formula for the heat transfer at low Mach number is derived in §2 and critically discussed in §3. §4 is devoted to its use on the problem at high Mach number which has just been described.

2. THE NEW METHOD IN THE INCOMPRESSIBLE CASE

In seeking an approximate method for the incompressible case $\rho = \text{constant}$ (for gases this refers to flow at low Mach number with frictional heating neglected), which will treat arbitrary main-stream velocity and wall-temperature distributions, one can choose between generalizing the methods (described in §1) of Fage & Falkner (1931) or of Squire (1942). The former seems more satisfactory since *a priori* restrictive approximations are applied therein only to the velocity distribution, not to the temperature distribution. Further, one finds that this approach simplifies greatly if the energy equation (2) is used in the von Mises form, when it becomes the second of equations (11) without the term representing frictional heating, namely

$$\frac{\partial T}{\partial x} = \frac{\mu\rho}{\sigma} \frac{\partial}{\partial \psi} \left(u \frac{\partial T}{\partial \psi} \right). \quad (17)$$

Following Fage & Falkner (1931), an approximation for u is inserted into equation (17) which is most closely accurate near the surface, namely $u = (\tau(x)/\mu)y$, where $\tau(x)$ is the skin friction at a distance x along the surface from the start of the boundary layer;† this implies

$$\psi = \rho \int_0^y u dy = (\rho\tau(x)/2\mu)y^2, \\ u = \left(\frac{2\tau(x)\psi}{\mu\rho} \right)^{\frac{1}{2}}. \quad (18)$$

and so

Then the equation becomes

$$\frac{\partial T}{\partial x} = \frac{1}{\sigma} (2\mu\rho\tau(x))^{\frac{1}{2}} \frac{\partial}{\partial \psi} \left(\psi^{\frac{1}{2}} \frac{\partial T}{\partial \psi} \right), \quad (19)$$

which will be solved under boundary conditions

(i) $T = 0$ at $x = 0$ and at $y = \infty$;‡

(ii) $T = T_0(x)$ and $k\partial T/\partial y = Q(x)$ at $y = 0$;

giving (equation (29) below) a relationship between $Q(x)$, $T_0(x)$ and $\tau(x)$.

The Heaviside operational method is used. If

$$t = \int_0^x \frac{1}{\sigma} (2\mu\rho\tau(z))^{\frac{1}{2}} dz, \quad (20)$$

and p is the Heaviside operator corresponding to $\partial/\partial t$, then, since $T = 0$ when $t = 0$, equation (19) becomes

$$pT = \frac{\partial}{\partial \psi} \left(\psi^{\frac{1}{2}} \frac{\partial T}{\partial \psi} \right). \quad (21)$$

Two linearly independent power series solutions of (21) are easily found, namely,

$$T = 1 + \frac{4p\psi^{\frac{1}{2}}}{2.3} + \frac{(4p\psi^{\frac{1}{2}})^2}{2.3.5.6} + \frac{(4p\psi^{\frac{1}{2}})^3}{2.3.5.6.8.9} + \dots \\ = \left(\frac{2}{3} p^{\frac{1}{2}} \psi^{\frac{1}{2}} \right)^{\frac{1}{2}} \left(-\frac{1}{3} \right)! I_{-\frac{1}{3}} \left(\frac{4}{3} p^{\frac{1}{2}} \psi^{\frac{1}{2}} \right), \quad (22)$$

† So that $\tau(x)/\mu = (\partial u/\partial y)_{y=0}$.

‡ Thus a temperature scale with the main-stream value as zero is used in this section.

and

$$T = \psi^{\frac{1}{2}} \left[1 + \frac{4p\psi^{\frac{1}{2}}}{3.4} + \frac{(4p\psi^{\frac{1}{2}})^2}{3.4.6.7} + \frac{(4p\psi^{\frac{1}{2}})^3}{3.4.6.7.9.10} + \dots \right]$$

$$= \psi^{\frac{1}{2}} \left(\frac{2}{3} p^{\frac{1}{2}} \psi^{\frac{1}{2}} \right)^{-\frac{1}{2}} \left(\frac{1}{3} \right)! I_{\frac{1}{2}} \left(\frac{4}{3} p^{\frac{1}{2}} \psi^{\frac{1}{2}} \right). \quad (23)$$

Now since, as $y \rightarrow 0$, by the boundary conditions (ii) above,

$$T = T_0(x) + \frac{Q(x)}{k} y + \dots = T_0(x) + \frac{Q(x)}{k} \left(\frac{2\mu}{\rho\tau(x)} \right)^{\frac{1}{2}} \psi^{\frac{1}{2}} + \dots, \quad (24)$$

the following combination of solutions (22) and (23) is required:

$$T = \left(\frac{2}{3} p^{\frac{1}{2}} \right)^{\frac{1}{2}} \psi^{\frac{1}{2}} \left(-\frac{1}{3} \right)! I_{-\frac{1}{2}} \left(\frac{4}{3} p^{\frac{1}{2}} \psi^{\frac{1}{2}} \right) T_0(x) + \left(\frac{2}{3} p^{\frac{1}{2}} \right)^{-\frac{1}{2}} \psi^{\frac{1}{2}} \left(\frac{1}{3} \right)! I_{\frac{1}{2}} \left(\frac{4}{3} p^{\frac{1}{2}} \psi^{\frac{1}{2}} \right) \frac{Q(x)}{k} \left(\frac{2\mu}{\rho\tau(x)} \right)^{\frac{1}{2}}. \quad (25)$$

(Here the functions of p 'operate on' the functions of x .)

The condition that $T \rightarrow 0$ as $\psi \rightarrow \infty$ must now be applied. Now the Bessel functions $I_{\lambda}(z)$, $I_{-\lambda}(z)$ tend to infinity (exponentially) as $z \rightarrow \infty$; and it is known that every linear combination of them also tends to infinity unless (like $K_{\lambda}(z)$) it is proportional to $I_{-\lambda}(z) - I_{\lambda}(z)$, when it tends exponentially to zero as $z \rightarrow \infty$. Therefore the coefficients of $I_{-\frac{1}{2}}$ and $I_{\frac{1}{2}}$ in (25) must be equal and opposite.

Hence

$$\frac{Q(x)}{k} \left(\frac{2\mu}{\rho\tau(x)} \right)^{\frac{1}{2}} = - \left(\frac{2}{3} \right)^{\frac{1}{2}} \frac{\left(-\frac{1}{3} \right)!}{\left(\frac{1}{3} \right)!} p^{\frac{1}{2}} T_0(x). \quad (26)$$

Now the operational form $p^{\frac{1}{2}} T_0(x)$ may be interpreted as

$$\frac{1}{\left(-\frac{1}{3} \right)!} \left[\frac{T_0(+0)}{t^{\frac{1}{2}}} + \int_0^t \frac{(dT_0/dt_1) dt_1}{(t-t_1)^{\frac{1}{2}}} \right] = \frac{1}{\left(-\frac{1}{3} \right)!} \int_0^x \frac{dT_0(x_1)}{(t-t_1)^{\frac{1}{2}}}, \quad (27)$$

where in the latter form (a 'Stieltjès integral'), which will be used for brevity's sake, there is a term in $T_0(+0)/t^{\frac{1}{2}}$ because the function $T_0(x)$ jumps from the value 0 which it has actually at the point $x = 0$, by condition (i) above, to a value which may be written $T_0(+0)$ immediately after $x = 0$. In (27) t_1 is the value of t at $x = x_1$.

By (20), expression (27) may be rewritten

$$\frac{(\sigma^2/2\mu\rho)^{\frac{1}{2}}}{\left(-\frac{1}{3} \right)!} \int_0^x \left(\int_{x_1}^x \sqrt{\{\tau(z)\}} dz \right)^{-\frac{1}{2}} dT_0(x_1), \quad (28)$$

which when substituted in equation (26) gives

$$Q(x) = -k \left(\frac{\sigma\rho}{9\mu^2} \right)^{\frac{1}{2}} \frac{\sqrt{\{\tau(x)\}}}{\left(\frac{1}{3} \right)!} \int_0^x \left(\int_{x_1}^x \sqrt{\{\tau(z)\}} dz \right)^{-\frac{1}{2}} dT_0(x_1). \quad (29)$$

This is the required relation between local heat transfer rate to the wall $Q(x)$, skin friction $\tau(x)$, and excess $T_0(x)$ of wall temperature over that of the stream.

The total heat transfer rate over an area of surface between $x = 0$ and $x = l$, per unit breadth, is therefore

$$\begin{aligned}\int_0^l Q(x) dx &= -\frac{k}{(\frac{1}{3})!} \left(\frac{\sigma\rho}{9\mu^2}\right)^{\frac{1}{3}} \int_0^l \sqrt{\{\tau(x)\}} dx \int_0^x \left(\int_{x_1}^x \sqrt{\{\tau(z)\}} dz\right)^{-\frac{1}{3}} dT_0(x_1) \\ &= -\frac{k}{(\frac{1}{3})!} \left(\frac{\sigma\rho}{9\mu^2}\right)^{\frac{1}{3}} \int_0^l dT_0(x_1) \int_{x_1}^l \sqrt{\{\tau(x)\}} \left(\int_{x_1}^x \sqrt{\{\tau(z)\}} dz\right)^{-\frac{1}{3}} dx \\ &= -\frac{\frac{1}{2}k}{(\frac{1}{3})!} \left(\frac{3\sigma\rho}{\mu^2}\right)^{\frac{1}{3}} \int_0^l \left(\int_x^l \sqrt{\{\tau(z)\}} dz\right)^{\frac{2}{3}} dT_0(x).\end{aligned}\quad (30)$$

Actually $(\frac{1}{3})! = 0.8930$.

3. CRITICAL DISCUSSION OF THE FORMULA. EXAMPLES

Clearly, the formula of §2 is in the nature of an asymptotic formula for large σ . For, when $\sigma = \nu/\kappa$ is large, the diffusivity of heat κ is much smaller than that of shear ν ; therefore the temperature boundary layer is only a small fraction of the velocity boundary layer, and in it the approximation $u = (\tau(x)/\mu)y$ is very good. Thus the work might be expected to be accurate merely for fluids like cold water ($\sigma \simeq 10$), oils ($\sigma > 100$), etc. However, it will actually be found to give good results even when σ is as low as 0.7 (which is the value, to one significant figure, for most gases including air, and is less than that of practically all liquids excepting liquid metals). To assess its accuracy more precisely, comparison with information from other sources, first qualitative and then quantitative, will now be made.

When all the physical constants of the fluid are altered, as well as the scale of the body and of the velocity distribution and of the distribution of difference between wall and stream temperatures, formula (30) tells us that a representative Nusselt number, $Nu = -\int_0^l Q(x) dx / kT_0$, where T_0 is an average value of $T_0(x)$, changes proportionately with $\sigma^{\frac{1}{3}}R^{\frac{1}{3}}$, where $R = Ul/\nu$ (with $u_1(x)$ proportional to U) is a representative Reynolds number. For it is well known that $\tau(x) \propto \rho U^2 R^{-\frac{1}{2}}$, and therefore

$$Nu \propto \left(\frac{\sigma\rho}{\mu^2} \rho U^2 R^{-\frac{1}{2}} l^2\right)^{\frac{1}{3}} = \sigma^{\frac{1}{3}} R^{\frac{1}{3}}. \quad (31)$$

The proportionality to $R^{\frac{1}{3}}$ is exactly true; it is almost obvious since, for given σ , the layer across which heat is being transferred has thickness $\propto R^{-\frac{1}{2}}$, and it is easily proved from the dynamical and thermodynamical equations. But the proportionality to $\sigma^{\frac{1}{3}}$ does not seem to have been previously suggested, as a general law exact for large σ and reasonably close even for σ of order 1.

In support of this aspect of the formula there are the following facts:

(i) $Nu R^{-\frac{1}{3}}$ is proportional to $\sigma^{\frac{1}{3}}$ to within 2% for a flat plate (Blasius layer) with uniform surface temperature, when $\sigma > 0.5$.

(ii) For the distribution $u = cx$ with uniform surface temperature, the calculated values of $Nu R^{-\frac{1}{3}}$ were interpolated (Goldstein 1938) as $0.57\sigma^{0.4}$ but fall below this value for large σ . Actually the slope of $\log(Nu R^{-\frac{1}{3}})$ against $\log \sigma$ is 0.39 for σ near 1 and 0.35 for σ near 10.

(iii) McAdams (1933) gives an empirical curve for heat transfer to a circular cylinder with laminar boundary layer,† from experiments with a wide range of liquids, with the equation $Nu = 0.86R^{0.43}\sigma^{0.3}$.

An alternative way of expressing (31) is to say that for a given body and given distributions of main-stream velocity and wall-stream temperature difference, the dependence of heat transfer upon the physical constants of the fluid is as

$$\text{Heat transfer} \propto k\sigma^{\frac{1}{2}}\nu^{-\frac{1}{2}} = k^{\frac{1}{2}}\rho^{\frac{1}{2}}c_p^{\frac{1}{2}}\mu^{-\frac{1}{2}}. \quad (32)$$

Yet another, in terms of the Péclet number $Pe = Ul/\kappa$, is as

$$Nu \propto Pe^{\frac{1}{2}}R^{\frac{1}{2}}, \quad (33)$$

showing that heat transfer is twice as sensitive to the thickness of the temperature boundary layer as to that of the velocity boundary layer.

Quantitative comparison is now sought between formula (30) and the results of more exact calculations. For most of these the wall temperature is uniform, and then equation (30) simplifies and may be written

$$Nu = 0.807 \left(\frac{\sigma\rho}{\mu^2} \right)^{\frac{1}{2}} \left(\int_0^l \sqrt{\{\tau(x)\}} dx \right)^{\frac{1}{2}}. \quad (34)$$

For the Falkner-Skan velocity distribution (1), $\tau(x) = \mu u_1 f''(0) (u_1/\nu x)^{\frac{1}{2}}$, and (34) becomes, with $R = u_1(l)/\nu$.

$$Nu R^{-\frac{1}{2}} = 0.807 \sigma^{\frac{1}{2}} [f''(0)]^{\frac{1}{2}} \left[\frac{3}{4}(m+1) \right]^{-\frac{1}{2}}. \quad (35)$$

The expression (35) is tabulated for $\sigma = 0.7$ in table 1, with values of $f''(0)$ taken from Hartree (1937). Values, presumably exact, calculated from equations (4) and (7) are adjoined. These were partly obtained from war-time German sources and partly computed by the present author.

TABLE 1

m	-0.0904	-0.0654	0	$\frac{1}{9}$	$\frac{1}{3}$	1	4	$\rightarrow \infty$
eqn. (35)	0	0.497	0.601	0.648	0.653	0.587	0.398	$\sim 0.921m^{-\frac{1}{2}}$
eqn. (7)	0.438	0.541	0.585	0.596	0.576	0.495	0.325	$\sim 0.741m^{-\frac{1}{2}}$
error (%)	-100	-8	+3	+9	+13	+18½	+22½	+24

The trend of the errors is easily understood when the shape of the velocity profiles for different m is considered and compared with the simple tangent at $y = 0$ with which each is replaced in the theory of § 2. Now the profile for $m = 0$ (that of Blasius) remains very close to its tangent near the surface, since for this profile both $\partial^2 u/\partial y^2$ and $\partial^3 u/\partial y^3$ vanish at $y = 0$. This explains why the agreement is best for $m = 0$; what there is of error is positive because the approximate u exceeds the accurate u (and the more rapid a flow the greater is the heat transfer by forced convection). For $m > 0$ the error is greater because the excess of $(\tau(x)/\mu)y$ over u is then much more marked (thus $\partial^2 u/\partial y^2 < 0$ for all y). For $m < 0$ the profile has a point of inflexion, $\partial^2 u/\partial y^2$ is positive for small y and negative for large y , and when this effect is at all

† This will include a certain amount of heat transferred through the wake.

marked the velocity in the crucial region near the surface is *underestimated* in using the tangent at $y = 0$, giving too small a heat transfer rate. This trend is accentuated to a ridiculous level for the 'separation' value $m = -0.0904$, for which $\tau(x)$ is identically zero.

In practical aerodynamics laminar boundary layers whose velocity profiles have a point of inflexion do not normally arise, since they become unstable at relatively low Reynolds numbers. On the other hand, the main-stream velocity does not normally increase more rapidly than cx , for some c . These facts limit the principal range of interest in the Falkner-Skan distribution to the interval $0 \leq m \leq 1$. In this interval the error in (35) varies from 2.7 to 18.6 %. If, however, the formula were multiplied by 0.904, the error in the modified formula would be limited to within ± 7.2 %. This would involve the substitution of 0.73 for 0.807 in (35).

Doubtless for large σ the formula is best used as it stands. But the considerations of the last paragraph indicate that for $\sigma = 0.7$, and generally (one supposes) for σ near 1, the formulae are improved if the constant $\frac{1}{2}3^{\frac{1}{2}}/(\frac{1}{3})! = 0.807$ be replaced by 0.73. The error should then normally be between -7.2 and $+7.2$ %, rising exceptionally to a maximum of $+12.3$ % for a main-stream velocity increasing like a high power of x , and falling to -17 % for the distribution $U = cx^{-0.0654}$, for which a marked point of inflexion already exists at $y = 2.05(\nu x/u_1)^{\frac{1}{2}}$. The main formula (30) becomes

$$\int_0^l Q(x) dx = -0.73\sigma^{\frac{1}{2}}k\nu^{-\frac{1}{2}} \int_0^l \left(\int_x^l (\tau(z)/\mu)^{\frac{1}{2}} dz \right)^{\frac{2}{3}} dT_0(x), \quad (36)$$

and here there is no reason to suppose that the errors will be increased by the non-uniformity of wall temperature since they are due simply to one cause (the use of the wrong velocity field). The simpler form for uniform wall temperature, corresponding to (34), can be written in terms of the skin friction coefficient $C_f = \tau(x)/(\frac{1}{2}\rho U^2)$ as

$$\text{Nu} = 0.73\sigma^{\frac{1}{2}}\nu^{-\frac{1}{2}} \left(\int_0^l (\tau(x)/\mu)^{\frac{1}{2}} dx \right)^{\frac{2}{3}} = 0.516\sigma^{\frac{1}{2}} (R \overline{\sqrt{C_f}})^{\frac{2}{3}}, \quad (37)$$

where the bar signifies a mean value along the surface $0 \leq x \leq l$. The latter form, a relation between Nu , σ , R and C_f , may be found the most useful of the present results. (The variation of the true exponent of σ , for values near $\sigma = 1$, from $\frac{1}{3}$ to a maximum of 0.39 for practical boundary layers, see (i) and (ii) above, hardly alters the above formula significantly; an alteration of $\sigma^{\frac{1}{2}}$ in (36) and (37) to $\sigma^{0.36}$ may be made if desired.)

On the other hand, when attention is confined to the boundary layer with uniform main stream (as will be the case in §4) the formula is best in its original form. Equation (29), with $\tau(x)$ replaced by $0.332\mu u_1 \sqrt{(u_1/\nu x)}$, its value for the Blasius profile, becomes

$$Q(x) = -0.339k\sigma^{\frac{1}{2}}(u_1/\nu)^{\frac{1}{2}} x^{-\frac{1}{2}} \int_0^x (x^{\frac{1}{2}} - x_1^{\frac{1}{2}})^{-\frac{2}{3}} dT_0(x_1), \quad (38)$$

with its integrated form (30) as

$$\int_0^l Q(x) dx = -0.677k\sigma^{\frac{1}{2}}(u_1/\nu)^{\frac{1}{2}} \int_0^l (l^{\frac{1}{2}} - x^{\frac{1}{2}})^{\frac{2}{3}} dT_0(x). \quad (39)$$

For the case when $T_0(x)$ is a series $\sum_0^\infty a_n x^n$, equation (38) may be compared with the exact value (8) calculated by Chapman & Rubesin (1949). It becomes (on putting $x_1 = x\xi$ inside the integrals)

$$\begin{aligned} Q(x) &= -0.339k\sigma^{\frac{1}{2}}(u_1/\nu x)^{\frac{1}{2}} \sum_0^\infty a_n x^n \int_0^1 (1-\xi^{\frac{2}{3}})^{-\frac{1}{2}} d(\xi^n) \\ &= -k(u_1/\nu x)^{\frac{1}{2}} \sum_0^\infty a_n x^n \left[0.339\sigma^{\frac{1}{2}} \frac{(-\frac{1}{3})! (\frac{4}{3}n)!}{(\frac{4}{3}n - \frac{1}{3})!} \right]. \end{aligned} \tag{40}$$

The quantity in square brackets is tabulated in table 2 against $-X'_n(0)$ for those values of n for which the latter were calculated by Chapman & Rubesin (1949), with $\sigma = 0.72$.

TABLE 2

n	0	1	2	3	4	5	10
$0.339 (0.72)^{\frac{1}{2}} (-\frac{1}{3})! (\frac{4}{3}n)! / (\frac{4}{3}n - \frac{1}{3})!$	0.304	0.490	0.594	0.671	0.734	0.787	0.987
$-X'_n(0) \dots$	0.296	0.489	0.597	0.684	0.744	0.799	1.006

For no value of n is the error as much as 3 %; yet the approximate values are calculated with the greatest simplicity. (Incidentally the fact that the error varies smoothly with n except for the value $n = 3$ indicates that a mistake must have been made in numerically solving the equation for $X_3(\eta)$; this fact can also be diagnosed by differencing the Chapman-Rubesin values of $X'_n(0)$. $X'_3(0)$ must really be about -0.678 ; thus in the Chapman-Rubesin notation $Y'_3(0) \simeq -1.356$.)

Hence equation (38) may be assumed to be accurate to within 3 %; its advantage not only in simplicity but in not requiring any series expansion for $T_0(x)$ will be apparent in §4, where it will be used to determine the surface-temperature distribution at high Mach number in the case when heat transfer to the wall is balanced by radiation from it.

4. ON COOLING BY RADIATION AT HIGH MACH NUMBER

As stated in §1, the only theory of any kind concerning the laminar boundary layer at high Mach number, under conditions of heat transfer, relates to the boundary layer with uniform main stream. For this the heat transfer rate is given (if the approximation $\mu \propto T$ be used) by the same equations at all Mach numbers. These are obtained by adding to the Pohlhausen solution (14) an arbitrary solution of (2) with $T = 0$ in the main stream, such as is given to within 3 %, as has just been shown, by

$$T = 0 \text{ in main stream, } T = T_2(x) \text{ at wall, } Q(x) = -0.339k\sigma^{\frac{1}{2}}\left(\frac{u_1}{\nu}\right)^{\frac{1}{2}} x^{-\frac{1}{2}} \int_0^x \frac{dT_2(x_1)}{(x^{\frac{2}{3}} - x_1^{\frac{2}{3}})^{\frac{1}{2}}}. \tag{41}$$

The sum may be written

$$Q(x) = 0.339k\sigma^{\frac{1}{2}}\left(\frac{u_1}{\nu}\right)^{\frac{1}{2}} \left\{ \left[T_1(1 + \tfrac{1}{2}(\gamma - 1)\sigma^{\frac{1}{2}}M^2) - T_0(+0) \right] x^{-\frac{1}{2}} - x^{-\frac{1}{2}} \int_0^x \frac{T'_0(x_1) dx_1}{(x^{\frac{2}{3}} - x_1^{\frac{2}{3}})^{\frac{1}{2}}} \right\}, \tag{42}$$

which shows how the solution (16) is modified for non-uniform wall-temperature distribution.

For a projectile the assumption of approximately uniform main stream, over the short front portion of the surface where the boundary layer is laminar, is probably quite good when the front shock wave is attached. This condition requires the head to be pointed and, for example, is satisfied by a wedge of semi-angle $< 45^\circ$ or a cone of semi-angle $< 60^\circ$ travelling symmetrically, point foremost, at high Mach number M . If the temperature of the atmosphere is Θ , and M' is the Mach number behind the shock, this latter is the Mach number which must be substituted in (42), while the temperature outside the boundary layer, T_1 , is given, by the condition of uniform stagnation temperature ('total temperature') everywhere outside the boundary layer (in particular across the shock), as

$$T_1 = \Theta \frac{1 + \frac{1}{2}(\gamma - 1) M^2}{1 + \frac{1}{2}(\gamma - 1) M'^2}. \quad (43)$$

Thus the 'theoretical final temperature' $T_f = T_1 [1 + \frac{1}{2}(\gamma - 1) \sigma^{\frac{1}{2}} M'^2]$, which the surface would attain if in its steady state there were no loss of heat from it by radiation or by conduction into the body, is slightly less than the stagnation temperature of the undisturbed stream, being in fact (for air, for which $\gamma - 1 \simeq \frac{2}{5}$, $\sigma^{\frac{1}{2}} \simeq \frac{5}{6}$)

$$T_f = T_1 [1 + \frac{1}{2}(\gamma - 1) \sigma^{\frac{1}{2}} M'^2] = \Theta (1 + \frac{1}{5} M^2) (1 + \frac{1}{6} M'^2) / (1 + \frac{1}{5} M'^2). \quad (44)$$

Now if heat transfer to the wall is entirely balanced by radiation from it (which will probably be its principal antagonist at high temperatures), then

$$Q(x) = \alpha [T_0(x)]^4, \quad (45)$$

where α is the Stefan-Boltzmann constant multiplied by the emissivity of the wall. The slight variation of α with temperature will be ignored, and in illustrative examples the emissivity will be taken as 70 % (a typical value for red-hot metals) so that $\alpha = 4 \times 10^{-5} \text{ ergs cm.}^{-2} \text{ sec.}^{-1} \text{ deg.}^{-4}$.

The only kind of estimate of the cooling effect of radiation which has hitherto been possible takes $T_0(x) = \text{constant}$ and averages the heat transfer balance along a length l of surface (say the length of the laminar layer). Thus it equates the total heat transfer rate in $0 < x < l$ to $\alpha T_0^4 l$. By (42) this gives

$$0.678 \sigma^{\frac{1}{2}} k (T_f - T_0) (u_1 l / \nu)^{\frac{1}{2}} = \alpha T_0^4 l, \quad (46)$$

or
$$\left(\frac{T_0}{T^*} \right)^4 + \frac{T_0}{T^*} = \frac{T_f}{T^*}, \quad (47)$$

where
$$T^{*3} = 0.678 \sigma^{\frac{1}{2}} k (u_1 l / \nu)^{\frac{1}{2}} / \alpha l. \quad (48)$$

The value of T^* for flight through the atmosphere depends on the altitude and on M and M' . Chapman & Rubesin (1949) point out that μ (and hence also k and ν) should be taken rather smaller than its main stream value, since then μ/T (which has been assumed constant) will be closer to its real values in the crucial region near the wall. To get a rough idea of the order of magnitude of T^* , the author took $M' = 2$, so that $u_1 = 2a$, where a is the velocity of sound behind the shock, and ignored the

change of $k(a/\nu)^{\frac{1}{2}}$ in passing through the shock, since, though $k\sqrt{a}$ increases by more than $\sqrt{\nu}$, a rather smaller value than that behind the shock (the 'main stream value') is in any case required. The physical quantities were given their values at 11 km. altitude in the 'standard atmosphere', namely, $\Theta = 217^\circ \text{K}$, $\sigma = 0.71$, $k = 2010 \text{ ergs cm.}^{-1} \text{sec.}^{-1} \text{deg.}^{-1}$, $a = 29500 \text{ cm. sec.}^{-1}$, $\nu = 0.293 \text{ cm.}^2 \text{sec.}^{-1}$, and l was taken as 10 cm. Then $T^* = 1620^\circ \text{K}$.

TABLE 3

T_f/T^*	0.5	1	1.5	2	2.5	3	3.5	4	4.5	5	5.5
T_0/T^*	0.453	0.724	0.885	1	1.09	1.16	1.23	1.29	1.33	1.38	1.42

Solutions to (47) are read off from table 3. Here, by (44), T_f is about $200(1 + \frac{1}{5}M^2)^\circ \text{K}$. For the illustrative value $T^* = 1620^\circ \text{K}$, table 4 gives the value of T_f and T_0 (in $^\circ \text{K}$)

TABLE 4

M	3	4	5	6	7	8	9	10	11	12	13	14	15
T_f	560	840	1200	1640	2160	2760	3440	4200	5040	5960	6960	8040	9200
T_0	535	760	980	1180	1360	1520	1660	1790	1910	2020	2130	2230	2330

against M . The cooling effect of radiation, on the above estimate which assumes it uniform and then averages its influence over the whole laminar layer, is seen to be very great. For example, the melting-points of iron (1810°K) and Fe_2O_3 (1830°K) are reached at a value of M exceeding 10, instead of one just over 6 as a theory neglecting radiation would indicate.

However, the indications of theory will be found less optimistic when the equation (46), which assumes $T_0(x)$ constant, is discarded, and $T_0(x)$ is determined by (more accurately) equating (42) to (45), giving

$$[T_0(x)]^4 = \frac{1}{2}l^{\frac{1}{2}}T^{*3} \left[\frac{T_f - T_0(+0)}{x^{\frac{1}{2}}} - \frac{1}{x^{\frac{1}{2}}} \int_0^x \frac{T'_0(x_1) dx_1}{(x^{\frac{1}{2}} - x_1^{\frac{1}{2}})^{\frac{1}{2}}} \right]. \quad (49)$$

Now the only possible value for $T_0(+0)$ in (49) is T_f . For if it had any other finite value, the integral must be of order $x^{-\frac{1}{2}}$ as $x \rightarrow 0$, and therefore $T'_0(x_1) = O(x_1^{-1})$ as $x_1 \rightarrow 0$, contradicting the finiteness of $T_0(+0)$. The physical reason is that heat transfer through the boundary layer is infinitely effective for x just greater than 0, since the layer is infinitesimally thin. Hence $T_0(+0) = T_f$, and $T_0(x)$ satisfies the equation

$$[T_0(x)]^4 = - \frac{T^{*3}l^{\frac{1}{2}}}{2x^{\frac{1}{2}}} \int_0^x \frac{T'_0(x_1) dx_1}{(x^{\frac{1}{2}} - x_1^{\frac{1}{2}})^{\frac{1}{2}}}. \quad (50)$$

The equation and initial condition may be rendered non-dimensional by putting

$$T_0(x) = T_f F \left[\left(\frac{T_f}{T^*} \right)^3 \left(\frac{x}{l} \right)^{\frac{1}{2}} \right], \quad (51)$$

when they become

$$[F(z)]^4 = - \frac{1}{2z^{\frac{1}{2}}} \int_0^z \frac{F'(z_1) dz_1}{(z^{\frac{1}{2}} - z_1^{\frac{1}{2}})^{\frac{1}{2}}}, \quad F(0) = 1. \quad (52)$$

A sufficiently accurate solution of the non-linear integral equation (52) was obtained as follows. Substitution of a series $F(z) = 1 - a_1 z + a_2 z^2 - a_3 z^3 + \dots$ in (52) yields that, with

$$A_n = 2[(\frac{2}{3}n - \frac{1}{3})!(-\frac{1}{3})!(-\frac{2}{3}n)!], \quad (54)$$

the a_n satisfy

$$(1 - a_1 z + a_2 z^2 - a_3 z^3 + \dots)^4 = (a_1/A_1) - (a_2/A_2)z + (a_3/A_3)z^2 - \dots, \quad (55)$$

whence they may be successively determined. In fact, one obtains

$$F(z) = 1 - 1.461z + 7.252z^2 - 46.46z^3 + 332.9z^4 - 2538z^5 + 20120z^6 - \dots \quad (56)$$

The ratio a_n/a_{n-1} is found to behave very regularly as a function of n^{-1} , and may be extrapolated to $n = \infty$ as 9.45. Hence the radius of convergence of (55) is the reciprocal of this, i.e. only 0.106. In fact $F(z)$ must have a singularity at $z = -0.106$. Actually it was investigated what possible algebraic singularity $F(z)$ could have at $z = -0.106$. The only possible behaviour was that $F(z) \sim 0.852(z + 0.106)^{-1/2}$, and that $F(z) - 0.852(z + 0.106)^{-1/2}$ is asymptotic to a multiple of $(z + 0.106)^{1/2}$, as $z \rightarrow -0.106$. If this multiple is determined approximately by the condition $F(0) = 1$, the resulting incipient series

$$F(z) = 0.852(z + 0.106)^{-1/2} - 0.198(z + 0.106)^{1/2} + \dots \quad (57)$$

might be assumed to have a much greater radius of convergence than (55), and therefore be used as a guide towards extrapolating (55) to the neighbourhood $z = 1$. (Actually it falls close to the curve determined below, being smaller by only 0.028 at $z = 1$.)

Next the behaviour of $F(z)$ for large z is obtained. A preliminary investigation (using a very general expansion) indicates an asymptotic series

$$F(z) \sim Az^{-1/2} + Bz^{-3/2} + \dots,$$

of which A and B are calculated as follows. Let w be a number between 0 and z , and potentially large. Then (52) may be written

$$\begin{aligned} (Az^{-1/2} + Bz^{-3/2} + \dots)^4 &= \frac{1}{2z^2} \left[- \int_0^\infty \frac{F'(z_1) dz_1}{(z^2 - z_1^2)^{3/2}} + \int_w^z \frac{(\frac{1}{2}Az_1^{-1/2} + \frac{1}{2}Bz_1^{-3/2} + O(z_1^{-5/2}))}{(z^2 - z_1^2)^{3/2}} dz_1 \right] \\ &= \frac{1}{2z} \left[- \int_0^w F'(z_1) dz_1 + O\left(\frac{w}{z}\right)^{1/2} \right] + \frac{1}{2} \int_{w/z}^1 \frac{\frac{1}{2}A(zu)^{-1/2} + \frac{1}{2}B(zu)^{-3/2} + O\{(zu)^{-5/2}\}}{(1 - u^2)^{3/2}} du. \end{aligned} \quad (58)$$

But, for large z/w ,

$$\begin{aligned} \int_{w/z}^1 u^{-1/2}(1 - u^2)^{-3/2} du &\sim 4(z/w)^{1/2}, & \int_{w/z}^1 u^{-3/2}(1 - u^2)^{-3/2} du &\sim 2(z/w)^{3/2}, \\ \int_{w/z}^1 u^{-5/2}(1 - u^2)^{-3/2} du &\sim \frac{4}{3}(z/w)^{5/2}. \end{aligned} \quad (59)$$

Further, it may be shown that, as $z/w \rightarrow \infty$,

$$\int_{w/z}^1 u^{-1/2}(1 - u^2)^{-3/2} du - 4\left(\frac{z}{w}\right)^{1/2} \rightarrow -2 \frac{(-\frac{1}{2})!(-\frac{1}{2})!}{(\frac{1}{2})!}. \quad (60)$$

ence, from (57), (58) and (59),

$$A^4 z^{-1} + 4A^3 B z^{-1} + O(z^{-1}) = \frac{1}{2z} \left[1 - A w^{-1} - B w^{-1} + O(w^{-1}) + O\left(\frac{w}{z}\right)^{\frac{1}{2}} \right] \\ + \frac{1}{2z} [A w^{-1} + B w^{-1} + O(w^{-1})] - \frac{1}{4} A z^{-1} \frac{(-\frac{1}{2})! (-\frac{1}{2})!}{(\frac{1}{2})!} + o(z^{-1} + z^{-1} w^{-1}). \quad (60)$$

By taking $w = z^{\frac{1}{2}}$, the O - and o -terms in (60) are all rendered $o(z^{-1})$, and therefore, equating coefficients of z^{-1} and z^{-1} , the formulae

$$A^4 = \frac{1}{2}, \quad 4A^3 B = -\frac{1}{4} A (-\frac{1}{2})! (-\frac{1}{2})! / (\frac{1}{2})! \quad (61)$$

are obtained, giving $A = 0.8409$, $B = -0.1524$. Actually the author obtained further terms in the asymptotic series for $F(z)$, by a similar process, namely,

$$F(z) \sim 0.8409 z^{-1} - 0.1524 z^{-1} - 0.0195 z^{-1} - 0.0038 z^{-1} + \dots \quad (62)$$

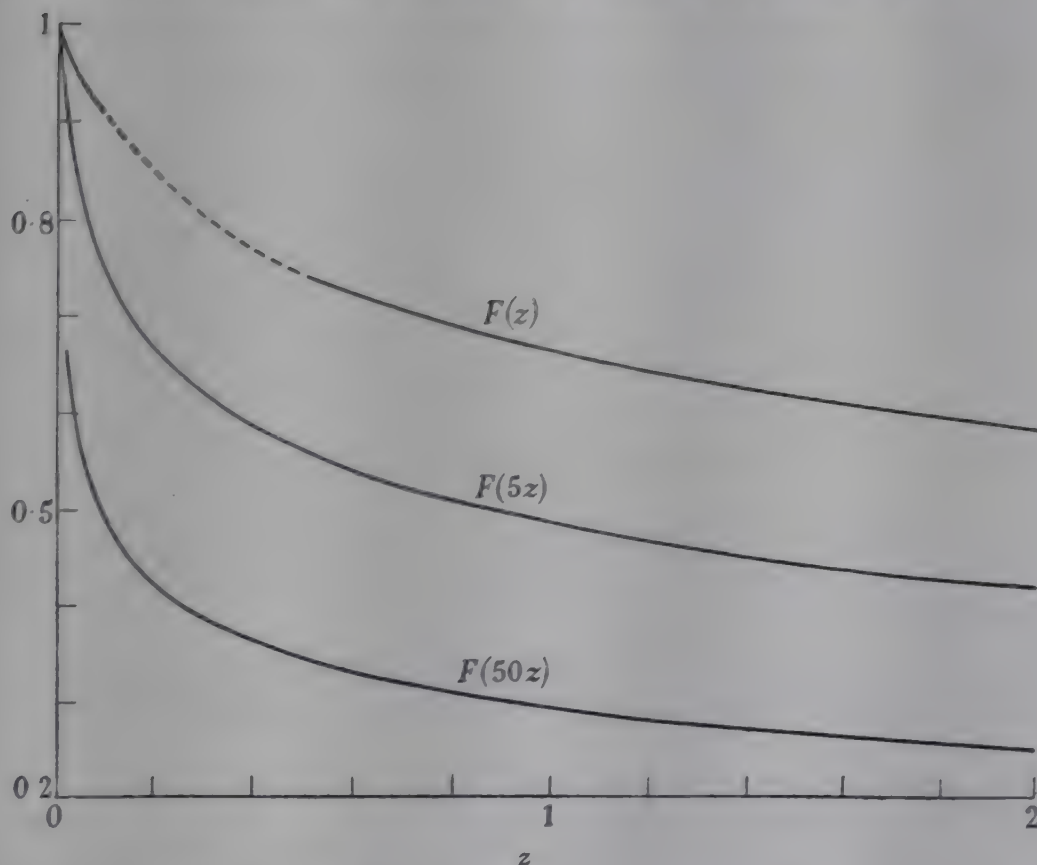


FIGURE 1. The function $F(z)$ which determines the surface temperature distribution by equation (51).

It seems likely (from the rate of drop of the coefficients) that the four terms displayed in (62) give sufficiently accurate values for $F(z)$ when $z > \frac{1}{2}$. Hence in figure 1 the series (55) for $F(z)$ is plotted as a plain curve in $0 < z < 0.1$, while (62) is plotted as a plain curve for $0.5 < z < 2.0$. They are simply faired together by the broken line in $0.1 < z < 0.5$, using (56) as a general guide to shape. This specification of $F(z)$ is then extended to $z = 100$ by plotting $F(5z)$ and $F(50z)$, in addition to $F(z)$, for $0 < z < 2$.

The function $F(z)$ which determines the wall-temperature distribution (51) being now known, the latter may be tabulated for $0 < x < l$ (i.e. all along the laminar layer) with various values of T_f/T^* . This is done in figure 2, where T_f/T^* is plotted against

x/l ($0 < x/l < 1$) for $T_f/T^* = 0.5, 1, 2$ and 4 . In each case T_f/T^* is the value of T/T^* at $x/l = 0$. (The reader is reminded that the reference temperature T^* was defined in (48), and that in an illustrative example T^* was 1620°K . With this value and with $T_f = 200(1 + \frac{1}{5}M^2)^\circ\text{K}$, the above values of T_f/T^* correspond to Mach numbers 3.9, 6.0, 8.7 and 12.5 respectively.) In each case comparison with the wall temperature deduced by assuming it constant and averaging the heat transfer balance has been made by nicking the curve at the point where T_0 takes this value (given by table 3).

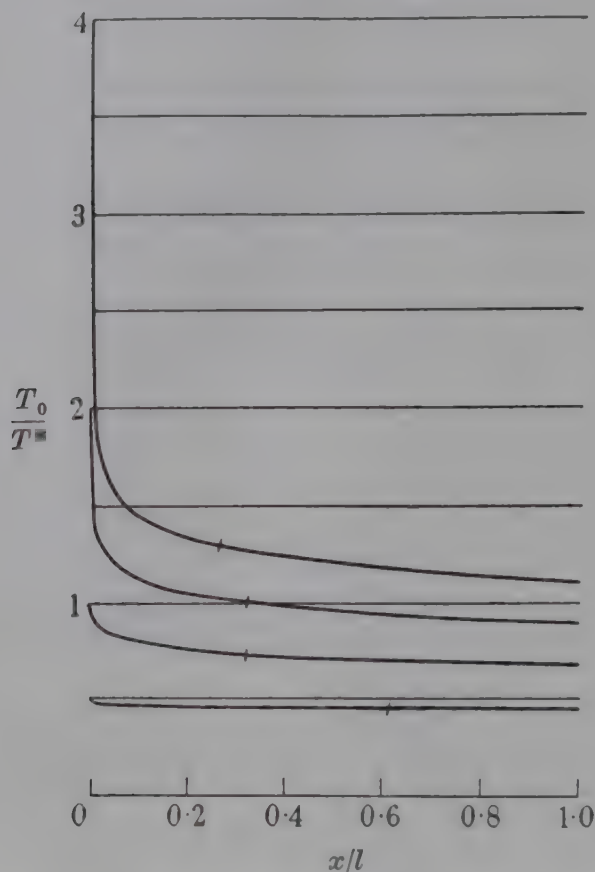


FIGURE 2. Theoretical distribution of surface temperature $T_0(x)$, relative to the reference temperature T^* , against proportionate distance x/l along the surface.

Even if the surface were thermally completely insulated the infinite temperature gradients where $x = 0$ and $T = T_f$ in the curves of figure 2 could not be credited. These are due to the assumption that the boundary-layer thickness is proportional to $x^{\frac{1}{2}}$, which ceases to be valid in the immediate neighbourhood of $x = 0$ where the boundary-layer equations are an inadequate approximation to the exact equations of gas dynamics. The boundary layer at $x = 0$ can be regarded as of finite thickness, of order ν/u_1 , indicating a local heat transfer rate of order $k(T_f - T_0(+0))u_1/\nu$, which when balanced with $\alpha[T_0(+0)]^4$ gives an equation of the form (47) for $T_0(+0)$ but with T^* replaced by

$$T^{**} = (ku_1/\nu\alpha)^{\frac{1}{4}}. \quad (63)$$

Since T^{**} is of order $30,000^\circ\text{K}$, this indicates that the peaks in the curves of figure 2 are rounded off but not far below $T = T_f$. However, conduction within the solid will far more effectively remove any high peak in the wall-temperature distribution. On the other hand, the more moderate variation from $x/l \simeq 0.2$ onwards will be altered

less considerably when conduction within the body is taken into account. The degree to which it is affected varies greatly with the material and shape of the body, but in each case the true distribution will presumably lie somewhere between the solution for zero body conductivity (as given in the curves of figure 2) and that for infinite body conductivity (which is roughly given by a flat straight line through the point nicked on each curve). In particular, the wall temperature will *fall from a maximum at the nose*. Thus bodies with sharp noses will at certain Mach numbers melt near the nose but not farther back.

It may be asked whether anything similar will occur when the front shock wave is detached, so that the boundary layer starts at a stagnation point, whence the main-stream velocity rises very rapidly to supersonic values. This question cannot be answered with any exactitude since no theory of the boundary-layer at high speeds, under heat transfer conditions, and with such a main-stream velocity distribution, exists. But if, as for the constant-pressure layer, the low-speed regime gives at least a rough indication of heat transfer rate, then it may be observed that, in the case considered above, the theory assuming wall temperature constant and averaging the heat-transfer balance underestimates the nose temperature simply because in fact the heat-transfer rate is largest at the nose when the wall temperature *is* constant, and therefore uniform radiation would be inadequate to keep the temperature down in this neighbourhood. Thus similar phenomena may be expected to occur with all main-stream velocity distributions for which $Q(x)$ is largest at $x = 0$ when $T_0(x) = \text{constant}$. This means that Nu increases less rapidly than l ; hence by (34) $\tau(x)$ increases less rapidly than x . This is known to be so for any main stream whose acceleration falls off from its initial value at the stagnation point, and will be most so for most rapid falling off of the acceleration.

Hence melting may occur (with possible resolidification further back) until the forward portion of the body is blunt enough to have a relatively low rate of fall of acceleration. Again this effect will be most marked for bodies with low thermal conductivity. Now the principal bodies with low thermal conductivity which move at a high enough Mach number to be melted by friction are stony meteorites. They enter the earth's atmosphere with a speed exceeding 10^4 m./sec., but usually with a considerable angular velocity as well. Thus only a small proportion of them can possess a well-defined 'nose' during flight, namely, those whose axis of rotation is in the direction of motion (for these the rotation will give their posture stability). Melting and vaporization will proceed rapidly at first, but finally (for those that reach the earth's surface at all) the Mach number will have been reduced to a value for which, on the above theory, melting should take place only near the nose.

There is some evidence that this occurs. Dr L. J. Spencer, F.R.S., in a letter to the author, says that, though only a few meteorites show a really well-oriented flight form, these few are of flat conical shape (sometimes with a rounded nose), with surface flow lines radiating from the nose and a ridge of fused material bordering the base. (The evidence for the fusion is petrological.) It is difficult to see how this can be explained except in terms of melting near the nose and resolidification farther back at a late stage of their flight.

APPENDIX

A method similar to that of § 2 can also be used to obtain an approximate integral equation for the skin friction $\tau(x)$ in terms of the main-stream velocity $u_1(x)$.

Von Mises's equation for the velocity field in an incompressible laminar layer is

$$\frac{\partial Z}{\partial x} = \mu \rho u \frac{\partial^2 Z}{\partial \psi^2}, \quad (64)$$

with ψ as in § 1 and $Z = u_1^2(x) - u^2$. The boundary conditions are $Z \rightarrow 0$ as $\psi \rightarrow \infty$, and as $x \rightarrow 0$, and

$$Z = u_1^2(x) - \left(\frac{\tau(x)}{\mu} y \right)^2 + O(y^3) = u_1^2 - \frac{2\tau(x)}{\mu \rho} \psi + O(\psi^{\frac{1}{2}}) \quad (65)$$

as $\psi \rightarrow 0$.

With the approximation (18) for u , equation (64) becomes

$$pZ = \psi^{\frac{1}{2}} \frac{\partial^2 Z}{\partial \psi^2}, \quad (66)$$

where p is the Heaviside operator corresponding to $\partial/\partial t$, where here

$$t = \int_0^x (2\mu \rho \tau(z))^{\frac{1}{2}} dz. \quad (67)$$

These equations are similar to (21) and (20); the solution of (66) satisfying (65) as $\psi \rightarrow 0$ is found as was that of (21) in § 2; it is

$$Z = \left(\frac{2}{3}p^{\frac{1}{2}}\right)^{\frac{1}{2}} \psi^{\frac{1}{2}} \left(-\frac{2}{3}\right)! I_{-\frac{1}{3}}\left(\frac{4}{3}p^{\frac{1}{2}}\psi^{\frac{1}{2}}\right) [u_1^2(x)] + \left(\frac{2}{3}p^{\frac{1}{2}}\right)^{-\frac{1}{2}} \psi^{\frac{1}{2}} \left(\frac{2}{3}\right)! I_{\frac{1}{3}}\left(\frac{4}{3}p^{\frac{1}{2}}\psi^{\frac{1}{2}}\right) \left[-\frac{2\tau(x)}{\mu \rho}\right]. \quad (68)$$

Since $Z \rightarrow 0$ as $\psi \rightarrow \infty$, the coefficients of $I_{-\frac{1}{3}}$ and $I_{\frac{1}{3}}$ in (68) must be equal and opposite. Hence

$$\begin{aligned} u_1^2(x) &= \frac{3^{\frac{1}{2}}\left(\frac{2}{3}\right)!}{2^{\frac{1}{2}}\left(-\frac{2}{3}\right)! \mu \rho} p^{-\frac{1}{2}} \tau(x) \\ &= \frac{3^{\frac{1}{2}}2^{\frac{1}{2}}}{\left(-\frac{2}{3}\right)! \mu \rho} \int_0^x \frac{\tau(x_1)}{(t-t_1)^{\frac{1}{2}}} \frac{dt_1}{dx_1} dx_1 \\ &= \frac{3^{\frac{1}{2}}2}{\left(-\frac{2}{3}\right)! \mu \rho} \int_0^x \left(\int_{x_1}^x \sqrt{\{\tau(z)\}} dz \right)^{-\frac{1}{2}} (\tau(x_1))^{\frac{1}{2}} dx_1. \end{aligned} \quad (69)$$

To estimate the accuracy of this integral equation for the skin friction, the solution when $u_1 = cx^m$ was obtained. This is given by

$$\tau(x) \left(\frac{x}{\mu \rho u_1^3} \right)^{\frac{1}{2}} = 3^{\frac{1}{2}} 2^{-\frac{1}{2}} (m+1)^{\frac{1}{2}} \left[\left(\frac{8m}{3m+3} \right)! \left(-\frac{2}{3} \right)! / \left(\frac{6m-2}{3m+3} \right)! \left(-\frac{1}{3} \right)! \right]^{\frac{1}{2}}. \quad (70)$$

The right-hand side is compared with the exact value (Hartree 1937) of the left-hand side ($f''(0)$ in the notation of § 1 above) in table 5.

TABLE 5

m	$-\frac{1}{9}$	-0.0904	-0.0654	0	$\frac{1}{9}$	$\frac{1}{3}$	1	4	$\rightarrow \infty$
eqn. (70)	0	0.101	0.175	0.312	0.473	0.698	1.136	2.218	$\sim 1.100 m^{\frac{1}{2}}$
Hartree (1937)	—	0	0.157	0.332	0.512	0.757	1.233	2.405	$\sim 1.193 m^{\frac{1}{2}}$
error (%)	—	$+\infty$	+11.5	-6	-7.6	-7.8	-7.9	-7.8	-7.8

It is seen that, in the whole range of m ($m \geq 0$) in which the main stream is not retarded and the velocity profile is therefore without point of inflexion, the error lies between -6 and -8% . Hence if the formula were multiplied by 1.075 the error would be less than 1% . This indicates that (69) thus revised, namely,

$$u_1^2(x) = 1.157 (\mu\rho)^{-\frac{1}{2}} \int_0^x \left(\int_{x_1}^x \sqrt{\{\tau(z)\}} dz \right)^{-\frac{1}{2}} (\tau(x_1))^{\frac{1}{2}} dx_1, \quad (71)$$

should give $\tau(x)$ correct to within 1% for most laminar layers occurring in practical aerodynamics. On the other hand, for layers with laminar separation it would get worse and worse as separation is approached; thus the separation value of m is given as -0.1111 instead of -0.0904 . Hence the use of (71) is no improvement on established approximate treatments of the laminar layer, such as Pohlhausen's, which are satisfactory except in regions of retarded flow. The integral equation will be less simple to solve than Pohlhausen's differential equation. Possibly (71) would be valuable, however, if the main-stream velocity which would produce a given skin friction distribution were required.

A solution of the integral equation (69) is possible explicitly if as an approximation the $\int_{x_1}^x \sqrt{\{\tau(z)\}} dz$ which is raised to the $(-\frac{1}{3})$ th power in the integrand is replaced by $(x-x_1) \sqrt{\{\tau(x_1)\}}$. Then the equation is invertible as

$$[\tau(x)]^{\frac{1}{2}} = \frac{(\mu\rho)^{\frac{1}{2}}}{(-\frac{1}{3})! 3^{\frac{1}{2}} 2} \int_0^x (x-x_1)^{-\frac{1}{3}} d[u_1^2(x)], \quad \tau(x) = 0.36 (\mu\rho)^{\frac{1}{2}} \left\{ \int_0^x (x-x_1)^{-\frac{1}{3}} d[u_1^2(x)] \right\}^{\frac{1}{2}}. \quad (72)$$

For $u_1 = cx^m$, this gives

$$\tau(x) \left(\frac{x}{\mu\rho u_1^3} \right)^{\frac{1}{2}} = 3^{-\frac{1}{2}} 2^{-\frac{1}{2}} \left(\frac{(2m)! (-\frac{2}{3})!}{(2m-\frac{2}{3})! (-\frac{1}{3})!} \right)^{\frac{1}{2}}, \quad (73)$$

which is tabulated in table 6. It is seen that the doubly approximate equation (72) can only be used, even in unretarded flow, when an error up to 10% can be tolerated. The errors near separation are still further increased.

TABLE 6

m	$-\frac{1}{6}$	-0.0904	-0.0654	0	$\frac{1}{6}$	$\frac{1}{3}$	1	4	$\rightarrow \infty$
eqn. (73)	0	0.214	0.260	0.360	0.495	0.698	1.112	2.158	$\sim 1.066 m^{\frac{1}{2}}$
error (%)	$-$	$+\infty$	$+65.6$	$+8.4$	-3.3	-7.8	-9.8	-10.3	-10.6

REFERENCES

- Chapman, D. R. & Rubesin, M. W. 1949 *J. Aero. Sci.* **16**, 547.
Crocco, L. 1946 *Monogr. Sci. Aeronaut.* no. 3.
Fage, A. & Falkner, V. M. 1931 *Rep. Memor. Aero. Res. Comm., Lond.*, no. 1408.
Falkner, V. M. & Skan, S. W. 1930 *Rep. Memor. Aero. Res. Comm., Lond.*, no. 1314.
Goldstein, S. (ed.) 1938 *Modern developments in fluid dynamics*, **2**, 631-632. Oxford Univ. Press.
Hartree, D. R. 1937 *Proc. Camb. Phil. Soc.* **33**, 223.
McAdams, W. H. 1933 *Heat transmission*, p. 224. New York: McGraw-Hill.
Pohlhausen, K. 1921 *Z. Angew. Math. Mech.* **1**, 120.
Squire, H. B. 1942 *Rep. Memor. Aero. Res. Comm., Lond.*, no. 1936.

The conductivity of thin wires in a magnetic field

BY R. G. CHAMBERS, *Royal Society Mond Laboratory, University of Cambridge*

(Communicated by Sir Lawrence Bragg, F.R.S.—Received 18 February 1950—
Revised 20 March 1950)

A thin film or wire of metal has a lower electrical conductivity than the bulk material if the thickness is comparable with or smaller than the electronic mean free path. Previous workers have obtained expressions for the magnitude of the effect by integrating the Boltzmann equation and imposing the appropriate boundary conditions. The problem is re-examined from a kinetic theory standpoint, and it is shown that the same expressions are obtained by this method, usually rather more simply, while the physical picture is considerably clarified. The method is applied to an evaluation of the conductivity of a thin wire with a magnetic field along the axis, and it is found that the resistivity should decrease as the magnetic field is increased; it should be possible to derive the mean free path and velocity of the conduction electrons by comparison of theory and experiment. The theory has been confirmed by experimental measurements on sodium; estimates of electronic velocity and mean free path are obtained which are in fair agreement with the values given by the free-electron theory.

INTRODUCTION

1. It has been realized for many years that the apparent resistivity of a conductor will be increased when its linear dimensions become comparable with, or smaller than, the electronic mean free path, and that from the magnitude of the effect the mean free path may be estimated. Thomson (1901) first gave an approximate expression for the increase in resistivity of a thin film; another approximation was given by Lovell (1936), and the exact solution for a free-electron metal was given by Fuchs (1938). For thin wires a very simple approximation due to Nordheim (1934) has been used by Eucken & Förster (1934) and succeeding workers. Interest in these problems has recently been revived by the work of Andrew (1949) on thin films of tin and thin wires of mercury, and particularly of MacDonald (1949), who showed that for a thin wire of sodium the resistivity actually decreased in a magnetic field, instead of increasing as for the bulk metal, because of the lengthening of the effective mean free path caused by the spiral motion of the electrons in the magnetic field. This opens up the possibility of determining not only the electronic mean free path, but also the momentum of the electrons at the surface of the Fermi distribution, since this determines the radius of the electronic orbits in a given field.

2. Sondheimer (1949) and MacDonald & Sarginson (1949) have recently given solutions for the variation in resistivity of a thin film with a magnetic field applied at right angles to the electric field, and either at right angles to or in the plane of the film. Experimentally, however, it is difficult to make reliable measurements on thin films, particularly on evaporated films, owing to the various possible disturbing effects (Reinders & Hamburger 1931; Lovell 1936–8), and experiments on thin wires may be easier to carry out. Dingle (1950) has recently given the exact theory of the increase in resistivity of a thin wire in the absence of a magnetic field, and MacDonald & Sarginson (in course of publication) have analyzed the case of a wire of rectangular,

and in particular of square, cross-section, and obtained results in the latter case similar to Dingle's.

The aims of the present paper are twofold: first, to show that an exact solution of these small-conductor problems can often be obtained very simply by kinetic theory arguments, without solving the Boltzmann equation *ab initio*, and secondly, to apply this method to the problem of a thin wire with a magnetic field along the axis. Experimental results are given which agree with the theoretical predictions.

THE KINETIC THEORY SOLUTION OF THIN-CONDUCTOR PROBLEMS

The case $H = 0$, $\epsilon = 0$

3. We consider first the case $H = 0$, $\epsilon = 0$, where ϵ is defined, following Fuchs, as the proportion of electrons which are specularly reflected at the surface of the metal, retaining their drift velocities, while the remainder $(1 - \epsilon)$ are diffusely reflected, with loss of drift velocity. The representation of the reflexion at the surface by a single parameter ϵ is, of course, very much of an idealization, just as is the representation of the scattering probability in the bulk metal by the single parameter l . Consider a point O in the metal and consider electrons passing through it in the direction $\mathbf{OP}$ where P is on the surface of the metal. Then assuming that the probability of an electron travelling a distance greater than x is $e^{-x/l}$, for $x < OP$, but that electrons which arrive at P will certainly collide there (so that $\epsilon = 0$), it is easily shown that the mean distance travelled by an electron without collision after passing through the point O is

$$l_1 = l(1 - e^{-OP/l}), \quad (1)$$

and it can similarly be shown that, for electrons travelling in the opposite direction $\mathbf{PO}$, the mean distance travelled without collision *before* reaching O is also given by equation (1).

In the presence of an electric field E_z in the z direction, these electrons will therefore have acquired a mean drift velocity

$$\Delta v_z = \frac{eE_z\tau}{mv} l(1 - e^{-OP/l}) = \frac{eE_z\tau}{m} (1 - e^{-OP/l}), \quad (2)$$

where v is the speed $|\mathbf{v}|$ of the electrons. We assume, in the usual way, that $N(\mathbf{v}, \mathbf{r})$, the electron density function, is only slightly different from $N_0(\mathbf{v}, \mathbf{r})$, and that $N_0(\mathbf{v}, \mathbf{r}) = N_0(|\mathbf{v}|)$ only. It is tempting at this point to suppose that the mean drift current density at the point O , due to electrons travelling in the direction $\mathbf{PO}$, is given simply by $eN_0(v)\Delta v_z$, and that the total drift current density at O is given by

$$j(O) = \frac{e^2 E_z \tau}{m} \int_0^\infty N_0(v) v^2 dv \int_0^{2\pi} d\phi \int_0^\pi d\theta \sin \theta (1 - e^{-OP/l}), \quad (3)$$

where the distance OP is to be expressed as a function of the angles θ and ϕ which define the direction of the vector $\mathbf{PO}$, θ being the angle between the direction $\mathbf{PO}$ and the z axis, and ϕ the azimuthal angle around the z axis. In fact, however, this approach neglects the fact that $N_0(|\mathbf{v}|)$ is not the same as $N_0(|\mathbf{v} + \Delta \mathbf{v}_z|)$; what

determines the drift current is the change n in the number of electrons travelling in the direction $\mathbf{PO}$ with velocity $\mathbf{v}$ due to the presence of the field. This is given by

$$n(\mathbf{PO}) = \frac{\partial N_0}{\partial v_z} \Delta v_z = \frac{eE_z \tau}{m} \frac{\partial N_0}{\partial v_z} (1 - e^{-OP/l}), \quad (4)$$

and the total drift current is given by

$$\begin{aligned} j(O) &= \int_0^\infty v^2 dv \int_0^{2\pi} d\phi \int_0^\pi d\theta \sin \theta e v_z n(\mathbf{PO}) \\ &= \frac{e^2 E_z \tau}{m} \int_0^\infty v^2 dv \int_0^{2\pi} d\phi \int_0^\pi d\theta \sin \theta v_z \frac{\partial N_0}{\partial v_z} (1 - e^{-OP/l}) \\ &= \frac{e^2 E_z \tau}{m} \int_0^\infty v^3 \frac{\partial N_0}{\partial v} dv \int_0^{2\pi} d\phi \int_0^\pi d\theta \sin \theta \cos^2 \theta (1 - e^{-OP/l}) \end{aligned} \quad (5)$$

When the distance OP is independent of (θ, ϕ) , the expression (5) reduces to expression (3), but not otherwise. It may be remarked that the slight difference noted by Reuter & Sondheimer (1948), between their formulation of the anomalous skin effect in metals at high frequencies and low temperatures and Pippard's (1947) formulation, consists in the absence of a $\cos^2 \theta$ term from the latter, and is due to the use of an expression of type (3) instead of type (5) by Pippard.

4. Equation (4) corresponds to the equation

$$n(\mathbf{r}, \mathbf{v}) = \frac{eE_z \tau}{m} \frac{\partial N_0}{\partial v_z} [1 - \exp\{-(\mathbf{r} - \mathbf{r}_0)/\tau \mathbf{v}\}] \quad [(\mathbf{r} - \mathbf{r}_0) \parallel \mathbf{v}],$$

which is a particular solution of the Boltzmann equation

$$\mathbf{v} \cdot \frac{\partial n}{\partial \mathbf{r}} - \frac{eE_z}{m} \frac{\partial N_0}{\partial v_z} = -\frac{n}{\tau},$$

if

$$\frac{\partial}{\partial \mathbf{r}} \cdot \left(\frac{eE_z}{m} \frac{\partial N_0}{\partial v_z} \right) = 0.$$

This is not, of course, the general solution of the Boltzmann equation; though it may be generalized to include the case $\epsilon \neq 0$, or to include the presence of a longitudinal magnetic field, the solution for a transverse magnetic field is not of this form—in fact, for a transverse field the kinetic theory equation (4) also has to be abandoned and a new approach adopted (§18). In this case the kinetic theory approach does become less easy to apply than the Boltzmann equation approach, though it can still shed useful light on the physical basis of the results obtained by the latter method.

5. We may at once write down the expression for the conductivity of a thin conductor by integrating (5) over the cross-sectional area s :

$$\sigma = \frac{e^2 \tau}{m} \int_0^\infty v^3 \frac{\partial N_0}{\partial v} dv \int_s ds \int_0^{2\pi} d\phi \int_0^\pi d\theta \sin \theta \cos^2 \theta (1 - e^{-OP/l})/s,$$

and if σ_0 is the bulk conductivity,

$$\sigma_0 = \frac{4\pi e^2 \tau}{3m} \int_0^\infty v^3 \frac{\partial N_0}{\partial v} dv,$$

so that

$$\begin{aligned}\frac{\sigma}{\sigma_0} &= \frac{3}{4\pi s} \int_s ds \int_0^{2\pi} d\phi \int_0^\pi d\theta \sin \theta \cos^2 \theta (1 - e^{-OP/l}) \\ &= 1 - \frac{3}{4\pi s} \int_s ds \int_0^{2\pi} d\phi \int_0^\pi d\theta \sin \theta \cos^2 \theta e^{-OP/l}.\end{aligned}\quad (6)$$

To find the value of σ/σ_0 for any shape of conductor it is merely necessary to express OP in terms of θ , ϕ and the position of the point O , and carry out the integrations. For a thin film or a thin wire, for instance, expressions are immediately obtained identical with those eventually found by Fuchs (1938, equation (17)) and Dingle (1950, equation (10.8)) by integration of the Boltzmann equation.

The case $H = 0$, $\epsilon \neq 0$

6. In the case $\epsilon \neq 0$ equation (1) has to be replaced by a slightly more complicated equation. For a thin film or a thin wire of circular cross-section (though not for arbitrary cross-section) it is easily seen that before travelling the distance OP , those electrons which have suffered repeated specular reflexions will have traversed equal distances PP' between successive reflexions from the wall, where PP' passes through O , is parallel to OP , and intersects the walls of the metal at P and P' . Letting $OP = a$, $PP' = b$, we find for the resultant mean free path

$$\begin{aligned}l_1 &= \int_0^a \frac{x}{l} e^{-x/l} dx + (1 - \epsilon) a e^{-a/l} + \epsilon \int_a^{a+b} \frac{x}{l} e^{-x/l} dx + \epsilon(1 - \epsilon)(a + b) e^{-(a+b)/l} \dots \\ &= l[1 - (1 - \epsilon)e^{-OP/l} / (1 - \epsilon e^{-PP'/l})],\end{aligned}\quad (7)$$

and equation (6) becomes

$$\frac{\sigma}{\sigma_0} = \frac{3}{4\pi s} \int_s ds \int_0^{2\pi} d\phi \int_0^\pi d\theta \sin \theta \cos^2 \theta [1 - (1 - \epsilon)e^{-OP/l} / (1 - \epsilon e^{-PP'/l})], \quad (8)$$

which, again, at once leads for a thin film or a thin wire to expressions identical with Fuchs's equation (21) and Dingle's equation (16.2) respectively. By expanding the denominators of these equations and comparing with the equations for $\epsilon = 0$, it may easily be shown that the following simple expression holds for both a thin wire and a thin film:

$$\left(\frac{\sigma}{\sigma_0}\right)_{\epsilon, \kappa} = (1 - \epsilon)^2 \sum_{n=1}^{\infty} \left\{ n \epsilon^{n-1} \left(\frac{\sigma}{\sigma_0}\right)_{\epsilon=0, n\kappa} \right\}, \quad (9)$$

which for a thin film may be verified by comparing Fuchs's equations (18) and (22). Here $\kappa = t/l$, where t is the thickness of the film or the diameter of the wire. It should be remarked that Fuchs's approximation (23) for small κ is incorrect; the correct expression is given in appendix 1 below.

The case $H \neq 0$. Longitudinal field

7. On the kinetic theory approach the case of a longitudinal magnetic field ($H \parallel E$) is easier to treat than the case of a transverse field, because in the first case the Lorentz force

$$\mathbf{F} = e(\mathbf{E} + \mathbf{v} \times \mathbf{H}/c) \quad (10)$$

splits up into two independent forces eE_z along the z axis and $e(\mathbf{v} \times H_z/c)$ which is always perpendicular to the z axis. We can then regard the electric field alone as producing a drift current in the normal way, and the magnetic field simply as modifying the electronic trajectories. When the electric and magnetic fields are not parallel, this approach has to be abandoned; this case is discussed briefly below (§18). For a longitudinal field and for $\epsilon = 0$, the equations (4), (5) and (6) are immediately applicable, where now OP is taken to be the distance from the point O to the point on the surface P along the trajectory of the electrons.

8. The problem of finding the conductivity in the presence of a magnetic field therefore reduces to the solution of equation (6). If the velocity of the electrons at the top of the Fermi distribution is v , then with a magnetic field H along the z axis, electrons travelling at an angle θ to the z axis will move in helical paths of which the projections on the (x, y) plane are circles of radius

$$r = \frac{mvc}{eH} \sin \theta = r_0 \sin \theta. \quad (11)$$

It should be noted that the value of r , and hence the effect of the field H , is independent of the direction of H , i.e. whether it is parallel or anti-parallel to E . If now, while the electron is travelling from P to O , its projection on the (x, y) plane traverses an angle ψ around such a circle, then the projection of the distance PO on the (x, y) plane is $\psi r_0 \sin \theta$, and the actual distance PO will be ψr_0 . Writing $l/r_0 = \eta$, therefore, equations (4) and (6) become

$$n = \frac{eE_z \tau}{m} \frac{\partial N_0}{\partial v_z} [1 - e^{-\psi/\eta}] \quad (12)$$

$$\text{and} \quad \frac{\sigma}{\sigma_0} = 1 - \frac{3}{4\pi s} \int_s ds \int_0^{2\pi} d\phi \int_0^\pi d\theta \cos^2 \theta \sin \theta e^{-\psi/\eta}, \quad (13)$$

where $\psi = \psi(x, y, \theta, \phi)$. The evaluation of this for a thin wire is discussed below.

In confirmation of equation (12), Dingle (unpublished) has found by direct integration of the Boltzmann equation for this case

$$n = \frac{eE_z \tau}{m} \frac{\partial N_0}{\partial v_z} \left[1 - \exp - \frac{1}{\eta} \left\{ \sin^{-1} \frac{\alpha r v_r}{[(v_r^2 + v_\theta^2)(v_r^2 + v_\theta^2 + 2\alpha r v_\theta + \frac{1}{2}\alpha r)]^{\frac{1}{2}}} \right. \right. \\ \left. \left. + \sin^{-1} \frac{\alpha [a^2(v_r^2 + v_\theta^2) - (r v_\theta - \frac{1}{2}\alpha a^2 - r^2)^2]^{\frac{1}{2}}}{[(v_r^2 + v_\theta^2)(v_r^2 + v_\theta^2 + 2\alpha r v_\theta + \frac{1}{2}\alpha r)]^{\frac{1}{2}}} \right\} \right], \quad (12a)$$

where $\alpha = eH/mc$, r is the distance of the point considered from the axis of the wire, and (v_r, v_θ, v_z) are components of velocity. It may be shown by a certain amount of trigonometry that the term $\{ \} = \psi$. This example shows rather clearly the advantage in clarity of the kinetic theory method.

Evaluation of equation (13) for a thin wire

9. The analytical solution of equation (13) is discussed in appendix 2: an explicit solution is possible only for large fields, and in general one must resort to numerical or graphical methods. By careful choice of procedure it is possible to obtain reasonably

accurate and detailed results without undue labour; the steps involved are described briefly below. As in the treatment of a thin wire in the absence of magnetic field, it is advantageous to integrate first over all electrons having the same value of ϕ , i.e. travelling in the same direction at the instant considered; the integral then becomes independent of ϕ by symmetry, and we may rewrite (13) in the form

$$\frac{\sigma}{\sigma_0} = \frac{3}{s} \int_0^{\frac{1}{2}\pi} d\theta \cos^2 \theta \sin \theta \int_s ds (1 - e^{-\psi/\eta})_\phi. \quad (14)$$

Now those electrons which, since colliding with the wall, have turned through angles between ψ and $\psi + d\psi$ and are, at the instant considered, travelling in the direction ϕ , θ will be located in a strip, of area $s(\psi) d\psi$ say, within the total cross-sectional area. If we denote the proportion of the total cross-section occupied by these electrons by $p(\psi) d\psi$, we may rewrite (14) as

$$\frac{\sigma}{\sigma_0} = 3 \int_0^{\frac{1}{2}\pi} d\theta \sin \theta \cos^2 \theta \int_0^\infty p(\psi) (1 - e^{-\psi/\eta}) d\psi, \quad (15)$$

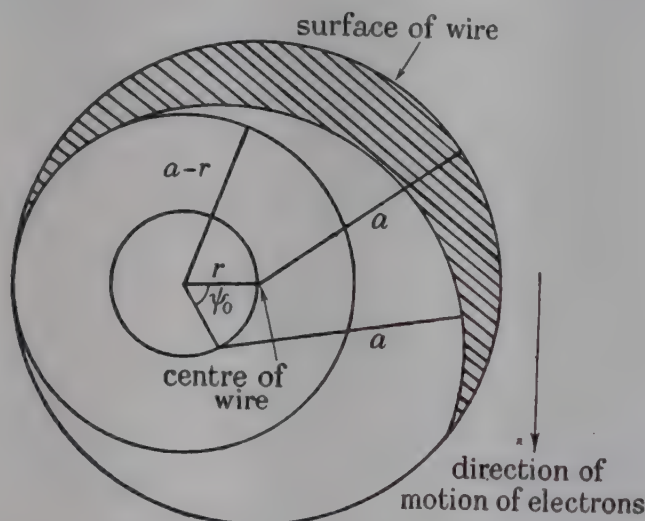


FIGURE 1. Construction for the determination of $P(\psi_0)$.

where ψ ranges between 0 and 2π for those electrons which can hit the wall at some point in their trajectory, and $\psi = \infty$ for those electrons which (in a strong enough field) can never hit the wall. Since for these $e^{-\psi/\eta} = 0$, we have

$$\frac{\sigma}{\sigma_0} = 3 \int_0^{\frac{1}{2}\pi} d\theta \sin \theta \cos^2 \theta \left[1 - \int_0^{2\pi} p(\psi) e^{-\psi/\eta} d\psi \right]. \quad (16)$$

Now $p(\psi)$ is a function only of the ratio of the radius of the wire, a , to the radius of the projection of the orbit, $r = r_0 \sin \theta$, i.e. of the ratio $2 \sin \theta / \eta \kappa$, where $\kappa = 2a/l$.

Geometrically, it is easier to find $P(\psi_0) = \int_0^{\psi_0} p(\psi) d\psi$ than to find $p(\psi)$ itself. Figure 1 shows the construction for the example $r/a = 2 \sin \theta / \eta \kappa = 0.3$, and $\psi_0 = 60^\circ$. If, at the instant considered, all electrons are travelling in the direction shown by the arrow, then the shaded area will contain all those which have turned through angles $\psi \leq \psi_0$, i.e. the area shaded is a measure of $P(\psi_0)$. The electrons within the circle of radius $(a - r)$ are those for which $\psi = \infty$. In this way $P(\psi_0)$ was evaluated for eight

values of r/a between 0.1 and 10. Now by partial integration, we may express the integral $\int_0^{2\pi} p(\psi) e^{-\psi/\eta} d\psi$ of equation (16) in terms of $P(\psi_0)$; we have

$$\int_0^{2\pi} p(\psi) e^{-\psi/\eta} d\psi = P(2\pi) e^{-2\pi/\eta} + \frac{1}{\eta} \int_0^{2\pi} P(\psi) e^{-\psi/\eta} d\psi. \quad (17)$$

Hence the values of $\int_0^{2\pi} p(\psi) e^{-\psi/\eta} d\psi$ were found, for the values of r/a considered, for values of η between 0.2 and 200. Denoting

$$1 - \int_0^{2\pi} p(\psi) e^{-\psi/\eta} d\psi \quad \text{by} \quad S_1(r/a, \eta) = S_1(2 \sin \theta / \eta \kappa, \eta) = S_1(\sin \theta / \kappa, \eta),$$

it was then possible to plot S_1 against $\kappa/\sin \theta$ for various values of η , as shown in figure 2. Finally, since $\sigma/\sigma_0 = \int_0^1 S_1 d(\cos^3 \theta)$, the values of σ/σ_0 for given κ and

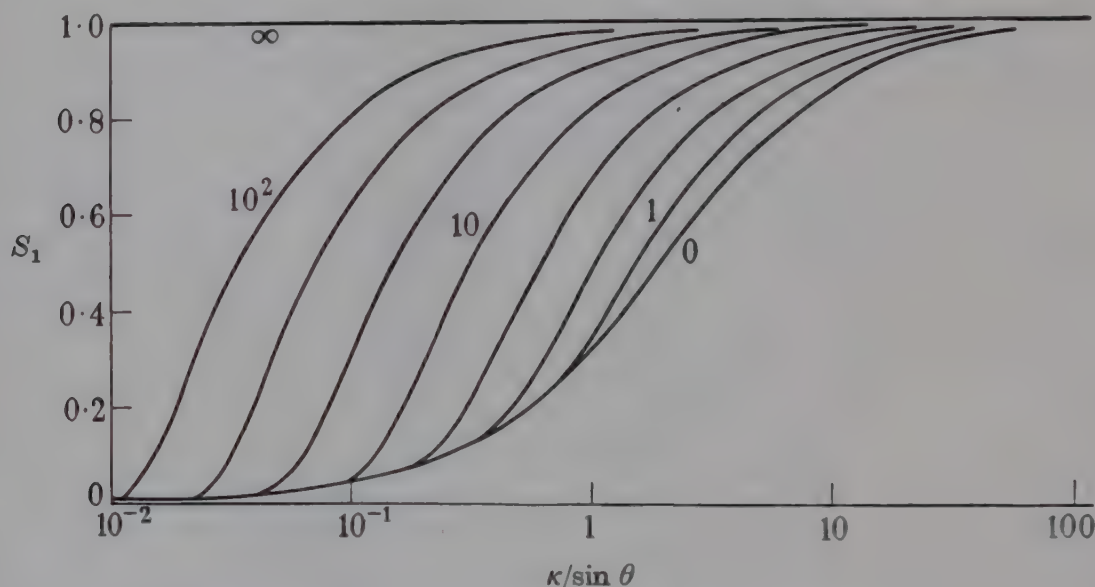


FIGURE 2. S_1 against $\kappa/\sin \theta$ for different values of $\frac{1}{2}\eta$.

η could be found by re-plotting S_1 against $\cos^3 \theta$, as in figure 3, plotted for $\kappa = 0.1$, $\eta = 0$ and $\eta = 20$, and integrating again. In this way σ/σ_0 was found for values of η between 0.2 and 200, and values of κ between 0.01 and 10. All integrations were performed graphically, with the aid of a planimeter. As a check on the accuracy, σ/σ_0 was at the same time evaluated by a very similar method for $\eta = 0$, i.e. zero magnetic field, and the results compared with the accurately known values (Dingle 1950).

10. The values found for σ/σ_0 are given in table 1. Here the values for $\eta = 0$ are those given by Dingle, and the values for $\eta\kappa \geq 5$ are those obtained from equation (A6) of appendix 2. The values given for the range $0 < \eta\kappa < 5$ are those obtained graphically as described above. Graphical solutions were also obtained for $\eta = 0$ and for $\eta\kappa \geq 5$; by comparison of the values so obtained with the exact values, and also by independent consideration of the inaccuracies involved, it is estimated that the values obtained graphically are probably not in error by more than 0.5 % for

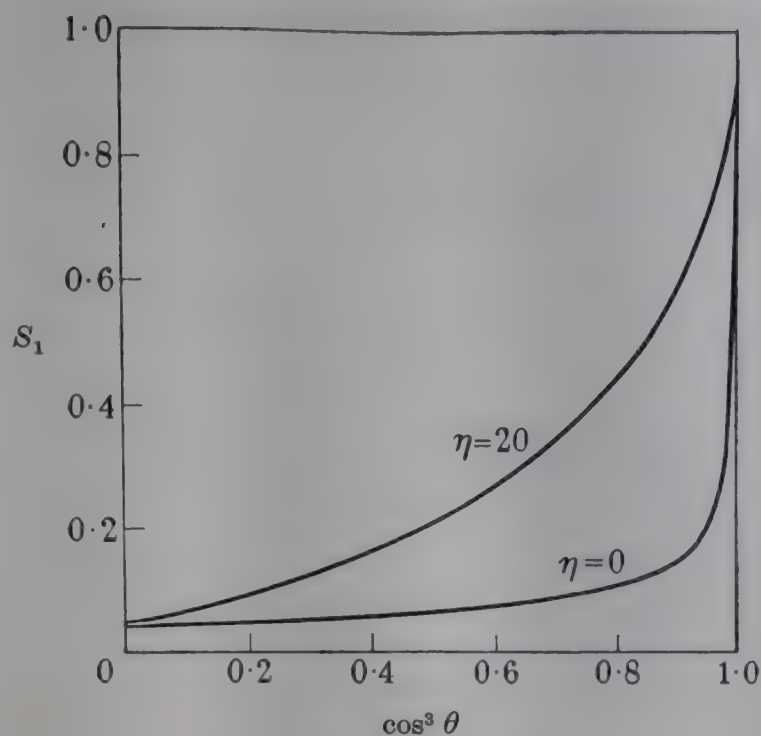


FIGURE 3. S_1 against $\cos^3 \theta$, for $\kappa = 0.1$, for $\eta = 0$ and $\eta = 20$.

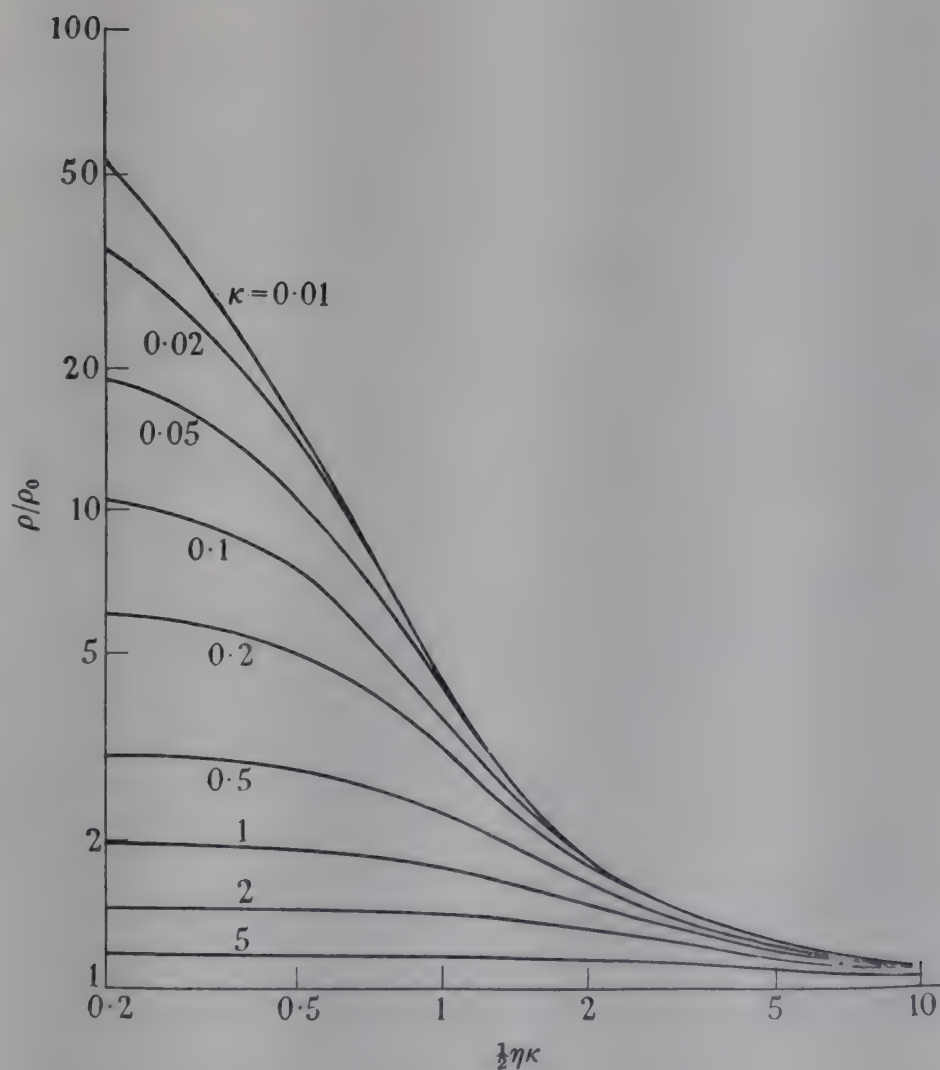


FIGURE 4. ρ/ρ_0 against $\frac{1}{2}\eta\kappa$, for various κ .

$\sigma/\sigma_0 \geq 0.5$, and more than 1 % for $\sigma/\sigma_0 \geq 0.1$; they are rather less accurate at smaller values, but seldom more than ± 3 % out even for the smallest values.

Figure 4 shows $\rho/\rho_0 = \sigma_0/\sigma$ plotted against $\frac{1}{2}\eta\kappa$, i.e. against $a/r_0 = eaH/mvc$. It will be seen that for large fields ρ/ρ_0 becomes dependent only on the product $\eta\kappa$, i.e. independent of l , as would be expected from the behaviour of equation (16) for large η .

TABLE 1. VALUES OF σ/σ_0 OBTAINED FROM EVALUATION OF EQUATION (13).

$\kappa \backslash \frac{1}{2}\eta$	0	10^{-1}	1	10^1	$10^{\frac{1}{2}}$	10	$10^{\frac{1}{2}}$	$10^{\frac{1}{2}}$	10^2
10	0.925	0.933	0.946	0.963	0.979	0.989	0.9948	0.9975	0.9988
5	0.853	0.869	0.893	0.928	0.959	0.979	0.9896	0.9950	0.9977
2	0.678	0.693	0.748	0.826	0.899	0.947	0.9741	0.9876	0.9942
1	0.489	0.504	0.555	0.674	0.805	0.896	0.9484	0.9753	0.9884
0.5	0.318	0.324	0.340	0.440	0.635	0.799	0.8984	0.9509	0.9768
0.2	0.158	—	0.163	0.186	0.295	0.549	0.758	0.8801	0.9425
0.1	0.0873	—	0.0899	0.0960	0.1266	0.272	0.5565	0.7691	0.8872
0.05	0.0464	—	—	0.0482	0.0556	0.0968	0.272	0.5746	0.7821
0.02	0.0192	—	—	—	0.0201	0.0281	0.0612	0.1992	0.5140
0.01	0.00976	—	—	—	—	0.0109	0.0200	0.0598	0.228

INTERPOLATION FORMULAE

11. For a thin wire in the absence of a magnetic field, the simple Nordheim expression $\sigma_0/\sigma \simeq 1 + 1/\kappa$ yields values of σ/σ_0 remarkably close to the true values (Dingle 1950). In the Nordheim theory, the resultant mean free path l_T is found from the equation $1/l_T = 1/l + 1/l_W$, where l is the bulk mean free path and l_W the mean distance to the walls. This is erroneous because the form of equation implies additive scattering probabilities, which in turn implies that the probability of making a collision with the wall varies exponentially with the distance travelled, which is untrue. Nevertheless, the differences between the values of σ/σ_0 given by the Nordheim equation and the true values are never greater than 5 % over the whole range of κ . This makes it very useful as an interpolation formula, when written in the form

$$\rho/\rho_0 = 1 + \alpha(\kappa)/\kappa, \quad (18)$$

where $\alpha(\kappa)$ is a function, found from the known values of $\rho/\rho_0(\kappa)$, which varies slowly from 1 to 0.75 as κ increases from 0 to ∞ . It was by the use of this function that the values of σ/σ_0 given in Dingle's paper for $\epsilon = \frac{1}{2}$ and for intermediate values of κ were found, using equation (9) above; this involved evaluating σ/σ_0 for $\epsilon = 0$ for a large number of values of κ , which was only made practicable by the use of equation (18). A similar simple interpolation function would be of considerable value in supplementing the values given in table 1. An attempt to find such a function by an adaptation of Nordheim's method was, however, unsuccessful; the resulting formulae were cumbersome and not particularly accurate. For $\eta\kappa \geq 2$, equation (A 6) of appendix 2 may be used; for $\eta\kappa \geq 10$ this equation may be simplified to

$$\frac{\sigma}{\sigma_0} \simeq 1 - \frac{3}{4\kappa(1 + \eta^2/4)} \left[1 + \frac{\eta^2}{8} (1 - e^{-2\pi/\eta}) \right] \quad (19)$$

with sufficient accuracy. Unfortunately, the region of large $\eta\kappa$ is the least interesting from the experimental point of view; it appears impossible to devise a simple and accurate formula valid in the important region $\eta\kappa \sim 1$.

APPLICATION TO EXPERIMENT

12. One of the major advantages of measurements using a magnetic field, compared with measurements made in zero field, is that all the required information may be obtained from one specimen only, by making measurements at various temperatures. This is in principle possible also with measurements in zero field, but only if the form of variation of σ_0 with temperature is accurately known. Since σ_0 cannot be measured on the specimen itself, this involves the doubtful assumption that the form of $\sigma_0(T)$, and therefore the form of $l(T)$, where l is the 'bulk' mean free path, is the same for a thin specimen as for a bulk specimen. Alternatively, σ may be measured at one temperature for a number of specimens of different diameter; this involves, however, the equally doubtful assumption that the 'bulk' mean free path is the same for all specimens, and that no variation in residual resistivity occurs between the specimens. It is probably some such variation which accounts for the rather large scatter of Andrew's (1949) observations about the theoretical curves (cf. Dingle 1950).

13. If we plot, not ρ/ρ_0 against $\frac{1}{2}\eta\kappa$ as in figure 4, but $\kappa\rho/\rho_0$ against $\frac{1}{2}\eta\kappa$, both the ordinates and abscissae become proportional to directly measurable quantities. We have $\kappa\rho/\rho_0 = 2a\rho/\rho_0 l$, and since $\rho_0 l$ is independent of temperature, we may write this $\rho/\rho_0 (l = 2a)$, the ratio of the observed resistivity at any temperature to the bulk resistivity at a temperature such that $l = 2a$. Also $\eta\kappa/2 = a/r_0 = (ea/mvc) H$, and is directly proportional to the applied magnetic field. If, therefore, we plot the observed values of ρ against H at a number of temperatures, on logarithmic scales, we shall obtain a family of curves which should be directly superposable on the theoretical curves of $\kappa\rho/\rho_0$ against $\eta\kappa/2$, and the proportionality constants between the two sets of curves will give $\rho_0 (l = 2a)$ and hence l , and (ea/mvc) and hence mv , directly.

14. In setting up and evaluating equation (13), we have tacitly made several assumptions which require further consideration before we can confidently compare the results with measurements on actual metals. First, we have assumed throughout that $\epsilon = 0$, i.e. that all electrons are diffusely reflected at the surface. Extension of the theory to $\epsilon \neq 0$ would be very laborious: equation (7) depends on symmetry considerations which are no longer applicable in the presence of a magnetic field, and it is unlikely that the simple result of equation (9) will still be valid. Moreover, the introduction of a third parameter ϵ besides κ and η , varying in an unknown way with temperature, would make it impossible to draw unique conclusions from the experimental results. For both these reasons, therefore, it is fortunate that independent evidence strongly indicates that $\epsilon = 0$ at all temperatures. This evidence comes from measurements by the author on the anomalous skin effect at high frequencies and low temperatures, where the mean free path becomes long compared with the classically predicted skin depth. The theory of the effect has been worked

out by Pippard (1947) and by Reuter & Sondheimer (1948) for the cases $p = 0$ and $p = 1$, where p is a parameter comparable with our ϵ . It has been found (Chambers 1950) that the observed behaviour of several metals agrees very closely with the $p = 0$ curve, and cannot be fitted to the $p = 1$ curve. It has been pointed out by Pippard that in the anomalous skin effect the current is carried almost entirely by electrons travelling at very small angles to the surface, and hence the value of p is presumably determined chiefly by the behaviour of these electrons on reflexion. If specular reflexion occurs at all, it seems most likely to occur for these electrons, travelling at small angles to the surface, and since it is not observed for them, i.e. since $p = 0$, it seems reasonable to expect that $\epsilon = 0$ also for the d.c. case, where electrons travelling at all angles to the surface are involved.

15. Secondly, we have assumed the validity of the classical kinetic theory or Boltzmann theory approach; that is to say, we have neglected quantization effects. Jones & Zener (1934) and Wilson (1936, p. 168) have pointed out that the Boltzmann equation ceases to be strictly applicable for $r_0 \leq l$, i.e. for $\eta \geq 1$ in our notation. However, it is usually assumed that quantization effects will not greatly alter the behaviour of the system unless the field is rather higher than those envisaged here.

16. Lastly, and most important, we have assumed a free-electron model, and we have assumed the existence of an isotropic mean path. It need hardly be said that these are both gross over-simplifications for any but group I metals; nevertheless, it might be expected that for more complex metals the experimental results would be of the same general form, and that from them some information about electronic mean free paths and momenta could be obtained. In one respect, however, the free-electron model predicts behaviour from which that of real metals departs strongly: it gives no bulk magneto-resistance effect, or at any rate an extremely small one. In fact, as Peierls (1931) first pointed out, bulk magneto-resistance effects are due entirely to the deviations of real metals from the free-electron model. This is confirmed by the experimental observations: the effect is very small in sodium, which most closely approximates to the free-electron model, and becomes progressively greater for less 'perfect' metals. Unfortunately, the effect is strongly temperature-dependent, becoming greater at low temperatures; in fact the relative change of resistance $\Delta\rho/\rho_0$ is approximately the same function of the ratio H/ρ_0 at all temperatures, so that as ρ_0 falls, the same relative increase in resistance will be brought about by a proportionately smaller field. (See, for example, Gruneisen (1945) for a review of the subject.) Thus the method described in § 13 to measure the mean free path and electron momentum can only be used for metals for which the bulk magneto-resistance effect is not too large; otherwise the increase of resistance with field due to this cause will swamp the decrease predicted for thin specimens by the free-electron theory. By comparison of the experimental data on the bulk effect with the expected behaviour of thin specimens, it appears that application of the method will probably only be possible for 'good' metals such as the group I metals, and perhaps not for all of them. Dr MacDonald has informed me that in measurements on thin films of gold, silver and tin, the bulk magneto-resistance effect was found to swamp the expected decrease in resistance in each case.

COMPARISON WITH EXPERIMENT

17. Through the kindness of Dr MacDonald, who lent me one of his specimens of sodium wire, 30μ in diameter, I have recently been able to make some experimental observations with which to compare the theoretical results given above. The specimen was placed in a cryostat capable of being heated electrically above the bath temperature, and the magnetic field was provided by a water-cooled solenoid capable of giving fields up to about 8000 gauss, uniform to within 1 or 2 % over the volume occupied by the specimen. The R - H curves obtained at various temperatures are

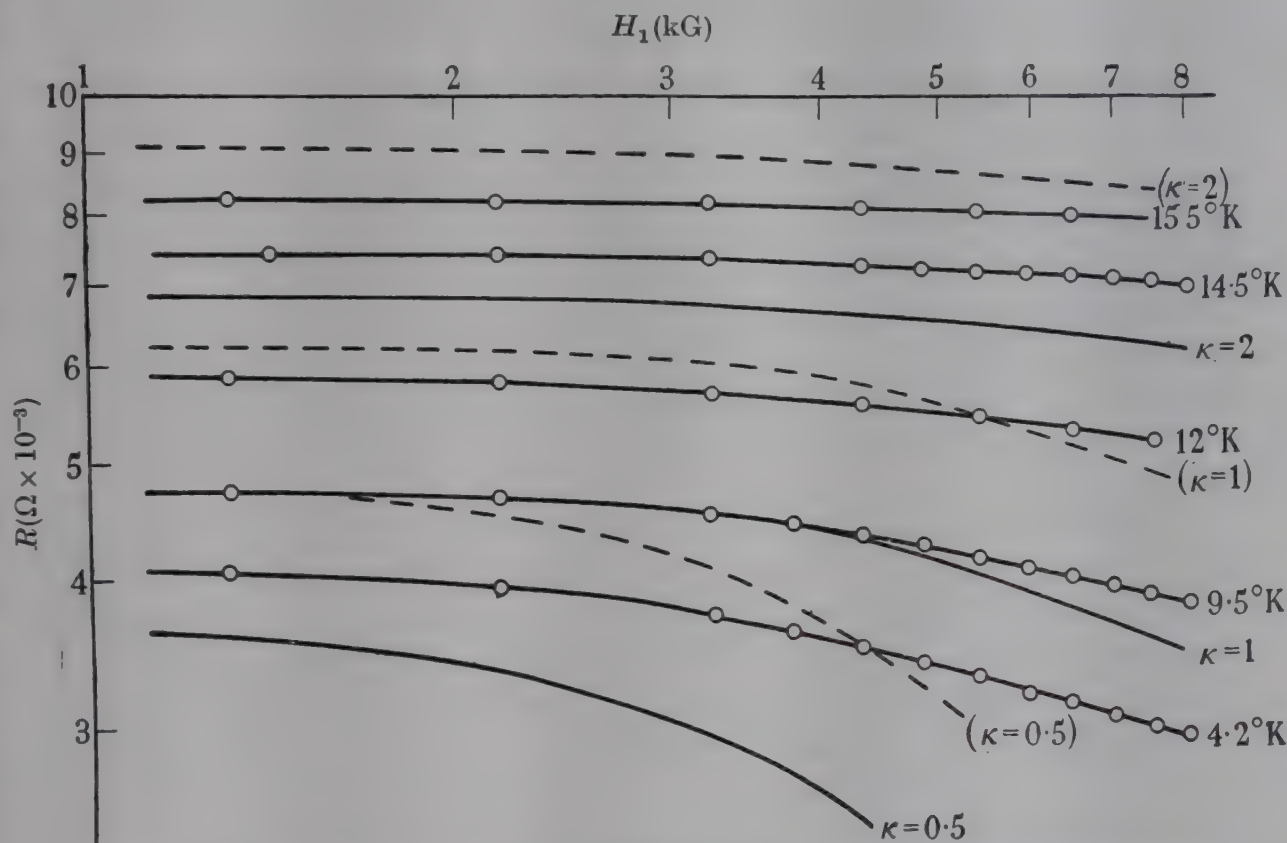


FIGURE 5. Experimental results on 30μ sodium wire, with theoretical curves.

shown in figure 5, and also shown are the theoretical curves for $\kappa = 0.5$, 1 and 2, adjusted to fit as well as possible. The departure from theory in high fields is presumably due to the bulk magneto-resistance effect discussed above. The magnitude of the departure is in good agreement with that expected from Justi's (1948) measurements on the bulk magneto-resistance effect in sodium. From the fit of the theoretical and experimental results, values for σ/l and for v of $9.0 \times 10^{10} \text{ ohm}^{-1} \text{ cm.}^{-2}$ and $1.0 \times 10^8 \text{ cm./sec.}$ are obtained, as described in § 13, which agree reasonably well with the values predicted by the free-electron theory, of $7.0 \times 10^{10} \text{ ohm}^{-1} \text{ cm.}^{-2}$ and $1.07 \times 10^8 \text{ cm./sec.}$ respectively.

The discrepancy between the theoretical and observed values of σ/l , however, appears to be real; the theoretical curves obtained assuming $\sigma/l = 7.0 \times 10^{10} \text{ ohm}^{-1} \text{ cm.}^{-2}$ are shown dotted in figure 5, and it will be seen that the discrepancies are too great to be accounted for by the bulk magneto-resistance effect.

TRANSVERSE FIELD

18. In conclusion, we may consider briefly the applicability of the kinetic theory method to transverse field problems. The equations of motion of a free electron subjected to electric and magnetic fields E_z and H_x are

$$\left. \begin{aligned} v_z &= v_0 \sin \left(\frac{eH_x}{mc} t + \delta \right), \\ v_y &= -v_0 \cos \left(\frac{eH_x}{mc} t + \delta \right) + \frac{cE_z}{H_x}, \end{aligned} \right\} \quad (20)$$

so that a freely moving electron acquires no mean drift velocity in the z direction, but simply 'precesses' in the y direction. If the electron is not freely moving, but suffers a collision after a time t , there will in general result currents in both the y and z directions; the mean current in the y direction is reduced to zero by the setting up of a Hall field E_y , and under these conditions the conductivity is given by $\bar{J}_z/E_z$. The existence of the Hall field greatly complicates the solution of these problems by kinetic methods for thin films and wires, particularly since, as pointed out by MacDonald & Sarginson, we can in general no longer assume E_y to be independent of y . In addition, the evaluation of the results for a case of such low symmetry as a thin wire in a transverse magnetic field would be extremely laborious; the relatively simple methods used in §9 for the longitudinal field case would no longer be applicable.

The particular case treated by Sondheimer can, however, be treated fairly simply, though not so elegantly as on the Boltzmann theory approach, and since in this case it is not obvious from simple physical arguments how the presence of the magnetic field affects the conductivity at all, we shall indicate briefly how this arises.

19. The case considered by Sondheimer is that of a thin film with an electric field E_z applied in the z direction and a magnetic field H_x perpendicular to the plane of the film. It is in this case possible to simplify the problem, since the Hall field is independent of position (the film being effectively of infinite extension in the (y, z) plane) by putting $E_y = 0$ and regarding E_z as the resultant of the applied field and the Hall field. Both $\bar{J}_z$ and $\bar{J}_y$ will then be non-zero; if the resultant total current is $\bar{\mathbf{J}}$, then the component of E_z in the direction of $\bar{\mathbf{J}}$ is to be regarded as the applied field, and the perpendicular component as the Hall field. Equations (20) then give the velocities v_y and v_z at time t of an electron following a path such that at time $t = 0$ it passes through a given point O in the metal with velocities

$$\left. \begin{aligned} v_z(0) &= v_0 \sin \delta, \\ v_y(0) &= -v_0 \cos \delta + \frac{cE_z}{H_x}. \end{aligned} \right\} \quad (21)$$

The velocity v_x of the electron in the x direction (across the film) is independent of time. Then if

$$v = [v_x^2 + v_y(0)^2 + v_z(0)^2]^{\frac{1}{2}}$$

is the speed of the electron at time $t = 0$, at time t its speed will, assuming $cE_z/H_x \ll v$, be given by

$$\Delta v_t = v_t - v = -\frac{v_0 c E_z}{v H_x} \left[\cos \left(\frac{eH_x}{mc} t + \delta \right) - \cos \delta \right]. \quad (22)$$

The excess number density of electrons having speed v_t , compared with those having speed v , is $(\partial N_0/\partial v) \Delta v_t = \delta n_t$ say, and these have a probability $e^{t/\tau}$ of reaching point O (t being negative).

Hence the excess number density of electrons travelling through point O with velocity components $v_x, v_y(0), v_z(0)$ is given by

$$n(0) = \int_{-T}^0 \frac{\partial}{\partial t} (\delta n_t) e^{t/\tau} dt = \frac{\partial N_0}{\partial v} \int_{-T}^0 \frac{\partial \Delta v_t}{\partial t} e^{t/\tau} dt, \quad (23)$$

where the limit of integration, $-T$, corresponds to electrons which have travelled from points at the surface of the metal. Using (22) and writing $t = x/v \cos \theta$, $T = x_0/v \cos \theta$, $r = mvc/eH_x$, $\tau = l/v$, we obtain from equation (23)

$$n(0) = \frac{\partial N_0}{\partial v} \frac{eE_z v_0}{m v v_x} \int_{-x_0}^0 \sin [(x/r \cos \theta) + \delta] e^{x/l \cos \theta} dx, \quad (24)$$

and on integrating this it becomes identical with the value of $n(0)$ given by Sondheimer, putting $E_y = 0$. The present solution makes clear the physical origin of the oscillations in conductivity found by Sondheimer. The oscillations in speed given by (22) give rise to oscillations in the excess number density at the points from which the electrons come which pass through the point O , the contribution to the current at O of electrons which have travelled freely from a point distance x away being given by the integrand of (24). It is clear that if l is large compared with x , and x is of the same order as r , the current density may fluctuate with the ratio x/r , as found by Sondheimer.

I should like to thank Dr MacDonald for the loan of the specimen of sodium wire on which the experimental measurements were made, and Dr MacDonald, Miss Sarginson, Dr Sondheimer and Mr Dingle for helpful discussions and for showing me their results before publication.

APPENDIX 1. *Explicit solution for thin films in zero magnetic field*

By the use of the approximation

$$\int_{\kappa}^{\infty} \frac{e^{-x}}{x} dx \simeq \ln 1/\kappa - c + 0.9755\kappa$$

(where c is Euler's constant), valid to $< 0.5\%$ for $\kappa < 0.2$, Fuchs's expression (18) for σ/σ_0 with $\epsilon = 0$ may be reduced to

$$\left(\frac{\sigma}{\sigma_0} \right)_{\epsilon=0, \kappa} \simeq \frac{3\kappa}{4} [\ln 1/\kappa + 0.4228] + 0.4816\kappa^2 \dots, \quad (A1)$$

valid to $< 1\%$ for $\kappa < 0.2$. Substituting this in equation (9), §6, we obtain for $\epsilon \neq 0$

$$\left(\frac{\sigma}{\sigma_0} \right)_{\epsilon, \kappa} \simeq \frac{1+\epsilon}{1-\epsilon} \left(\frac{\sigma}{\sigma_0} \right)_{\epsilon=0, \kappa} + \frac{4\epsilon}{(1-\epsilon)^2} 0.4816\kappa^2 - \frac{3\kappa}{4} (1-\epsilon)^2 \sum_1^{\infty} n^2 \epsilon^{n-1} \ln n. \quad (A2)$$

This expression is applicable provided that the terms in the summation (9) for which $n\kappa \geq 0.2$ are small compared with those for which $n\kappa < 0.2$. For $\epsilon = \frac{1}{2}$, this restricts

its application to values of κ less than about 0.02, and for $\epsilon = \frac{9}{10}$, less than about 0.004. The term $\sum_1^\infty n^2 \epsilon^{n-1} \ln n$ in (A 2) is equal to 15.48 for $\epsilon = \frac{1}{2}$, and 6060 for $\epsilon = \frac{9}{10}$.

This expression gives values of $(\sigma/\sigma_0)_{\epsilon, \kappa}$ somewhat greater than those given by Fuchs's equation (23) (in which $1 - \epsilon$ is presumably a misprint for $(1 - \epsilon)^{-1}$) and plotted in his figure 2, but the difference does not materially affect his interpretation of Lovell's experimental results.

APPENDIX 2. Analytical solution of equation (13)

1. The analytical approach to the solution of equation (13) is identical with the graphical approach of §9, and the same notation will be used. Let the ratio $r/a = \gamma$. Then by consideration of figure 1, and the corresponding figure for $\gamma > 1$, it is readily shown that for $\gamma \leq 1$,

$$\pi p(\psi) = \gamma(1 - \gamma) + \gamma^2 \cos^2 \frac{1}{2}\psi + \gamma \cos \frac{1}{2}\psi (1 - \gamma^2 \sin^2 \frac{1}{2}\psi)^{\frac{1}{2}} \quad (0 \leq \psi \leq 2\pi) \quad (\text{A } 3)$$

and for $\gamma \geq 1$,

$$\left. \begin{aligned} \pi p(\psi) &= 2\gamma \cos \frac{1}{2}\psi (1 - \gamma^2 \sin^2 \frac{1}{2}\psi)^{\frac{1}{2}} & (0 \leq \psi \leq 2 \sin^{-1}(1/\gamma)) \\ &= 0 & (2 \sin^{-1}(1/\gamma) < \psi \leq 2\pi). \end{aligned} \right\} \quad (\text{A } 4)$$

2. Now for $\eta\kappa \geq 2$, $\gamma = r/a = 2 \sin \theta / \eta\kappa$ remains less than 1 for all values of θ , while for $\eta\kappa < 2$, we have $\gamma < 1$ for small θ and $\gamma > 1$ for large θ ($\sim \frac{1}{2}\pi$). Physically, this means simply that for strong enough fields ($\eta\kappa \geq 2$) all electronic trajectories are curved into paths of radius less than the wire radius; for smaller fields, those electrons travelling at small angles to the axis and therefore having small transverse velocities will still follow such paths, but electrons moving at greater angles to the axis will follow paths of radius greater than the wire radius. Solution for $\eta\kappa < 2$ is therefore much more difficult than for $\eta\kappa \geq 2$, and will only be considered briefly.

3. The value of $S_1 = 1 - \int_0^{2\pi} p(\psi) e^{-\psi/\eta} d\psi$ may be found by expanding the root $(1 - \gamma^2 \sin^2 \frac{1}{2}\psi)^{\frac{1}{2}}$ in (A 3) and (A 4). For $\gamma < 1$, the series solution converges rapidly; taking

$$(1 - \gamma^2 \sin^2 \frac{1}{2}\psi)^{\frac{1}{2}} \simeq 1 - \frac{1}{2}\gamma^2 + \frac{1}{2}\gamma^2 \cos^2 \frac{1}{2}\psi,$$

we obtain

$$\begin{aligned} \pi(1 - S_1) &\simeq \gamma \left[\frac{2}{\alpha} (1 - e^{-\pi\alpha}) + \frac{2\alpha}{\alpha^2 + 1} (1 + e^{-\pi\alpha}) \right] + \gamma^2 \left[\frac{\alpha}{\alpha^2 + 4} - \frac{1}{\alpha} \right] (1 - e^{-\pi\alpha}) \\ &\quad + \frac{1}{4}\gamma^3 \left[\frac{\alpha}{\alpha^2 + 9} - \frac{\alpha}{\alpha^2 + 1} \right] (1 + e^{-\pi\alpha}) \quad (\text{A } 5) \end{aligned}$$

where $\alpha = 2/\eta$. Hence, for $\eta\kappa > 2$, so that $\gamma < 1$ for all θ , we have

$$\begin{aligned} \frac{\sigma}{\sigma_0} &= 3 \int_0^{\frac{1}{2}\pi} S_1 \cos^2 \theta \sin \theta d\theta \\ &\simeq 1 - \frac{3}{4\kappa(1 + \frac{1}{4}\eta^2)} \left[1 + \frac{\eta^2}{8} (1 - e^{-2\pi/\eta}) \right] + \frac{3}{16\kappa^2} \frac{1}{(1 + \frac{1}{4}\eta^2)(1 + 9\eta^2/4)} (1 + e^{-2\pi/\eta}) \\ &\quad + \frac{4}{5\pi\kappa^2} \frac{\eta}{1 + \eta^2} (1 - e^{-2\pi/\eta}) \quad (\text{A } 6) \end{aligned}$$

This expression was used in deriving the values of σ/σ_0 for $\eta\kappa \geq 5$ given in table 1. The series converges very rapidly; the next two terms, given by the next two terms in the expansion of $(1 - \gamma^2 \sin^2 \frac{1}{2}\psi)^{\frac{1}{2}}$, give a contribution less than

$$\frac{0.1\eta}{1 + \frac{1}{4}\eta^2} [(\eta\kappa)^{-5} + 1.6(\eta\kappa)^{-7}],$$

and are consequently always negligible for $\eta\kappa \geq 5$, and only become important for $\eta\kappa \sim 2$ when $\eta \sim 2$.

4. For $\eta\kappa < 2$, an expression for S_1 may still be obtained by expanding the root in (A3) and (A4), but for $\gamma > 1$ the resulting expression converges only slowly, because the value of $(1 - \gamma^2 \sin^2 \frac{1}{2}\psi)^{\frac{1}{2}}$ falls to zero at the upper limit of integration. Moreover, to obtain the value of σ/σ_0 from S_1 it is necessary to evaluate integrals of the form

$$\int_{\sin^{-1}(\frac{1}{2}\eta\kappa)}^{\frac{1}{2}\pi} d\theta \sin^{2n} \theta \exp \{ -\alpha \sin^{-1} (\eta\kappa/2 \sin \theta) \}.$$

Integration of this has not been attempted. It may be noted, however, that for very small fields, i.e. $\eta\kappa \rightarrow 0$, where we may replace (A4) by $\pi p(\psi) \simeq 2\gamma(1 - \frac{1}{4}\gamma^2\psi^2)^{\frac{1}{2}}$, and where we may take $\gamma > 1$ over the whole range of θ from 0 to $\frac{1}{2}\pi$, the expression for σ/σ_0 reduces to

$$\frac{\sigma}{\sigma_0} \simeq 1 - \frac{12}{\pi} \int_0^1 dx (1 - x^2)^{\frac{1}{2}} \int_0^{\frac{1}{2}\pi} d\theta \cos^2 \theta \sin \theta e^{-\kappa x / \sin \theta},$$

where we have put $\gamma\psi = 2x$. This is identical with Dingle's equation (10.8) for a thin wire in the absence of a magnetic field. It would be possible by going to higher approximations to find an explicit solution for σ/σ_0 in very small fields, in terms of the value of σ/σ_0 for $H = 0$, but this could only be valid for $\eta\kappa \leq 0.1$, and reference to table 1 and figure 4 shows that in this region σ/σ_0 changes only very slightly.

REFERENCES

- Andrew, E. R. 1949 *Proc. Phys. Soc.* **62**, 77.
 Chambers, R. G. 1950 *Nature*, **165**, 239.
 Dingle, R. B. 1950 *Proc. Roy. Soc. A*, **201**, 545.
 Eucken, A. & Förster, F. 1934 *Nach. Ges. Wiss. Göttingen*, N.S. 1-2, **43**, 129. See also *Ann. Phys., Lpz.*, **28**, 603 (1937), **33**, 733 (1938) (Riedel); **30**, 494 (1937) (Reuter); **40**, 121 (1941) (Möser).
 Fuchs, K. 1938 *Proc. Camb. Phil. Soc.* **34**, 100.
 Gruneisen, E. 1945 *Ergebn. Exakt. Naturw.* **21**, 50.
 Jones, H. & Zener, C. 1934 *Proc. Roy. Soc. A*, **144**, 101.
 Justi, E. 1948 *Ann. Phys., Lpz.* (6) **3**, 183.
 Lovell, A. C. B. 1936 *Proc. Roy. Soc. A*, **157**, 311.
 Lovell, A. C. B. 1938 *Proc. Roy. Soc. A*, **166**, 270.
 Lovell, A. C. B. & Appleyard, E. T. S. 1937 *Proc. Roy. Soc. A*, **158**, 718. See also Appleyard, E. T. S. & Bristow, J. R. 1939, *Proc. Roy. Soc. A*, **172**, 530.
 MacDonald, D. K. C. 1949 *Nature*, **163**, 639.
 MacDonald, D. K. C. & Sarginson, K. 1949 *Nature*, **164**, 921.
 Nordheim, L. 1934 *Act. Sci. et Ind.* no. 131. Paris: Hermann. See also Planck, W. 1914, *Phys. Z.* **15**, 569; Tisza, L. 1931, *Naturwissenschaften*, **19**, 86; Nordheim, L. 1931, *Ann. Phys., Lpz.*, **9**, 607.

- Peierls, R. 1931 *Ann. Phys., Lpz.*, **10**, 97.
 Pippard, A. B. 1947 *Proc. Roy. Soc. A*, **191**, 385.
 Reinders, W. & Hamburger, L. 1931 *Ann. Phys., Lpz.*, **10**, 649, 789.
 Reuter, G. E. H. & Sondheimer, E. H. 1948 *Nature*, **161**, 394; *Proc. Roy. Soc. A*, **195**, 336.
 Sondheimer, E. H. 1949 *Nature*, **164**, 920.
 Thomson, J. J. 1901 *Proc. Camb. Phil. Soc.* **11**, 120.
 Wilson, A. H. 1936 *Theory of metals*. Cambridge University Press.

The structure of linear relativistic wave equations.

II. Representations

BY K. J. LE COUTEUR

Department of Theoretical Physics, University of Liverpool

(Communicated by P. A. M. Dirac, F.R.S.—Received 22 February 1950)

A representation of the spin matrices $I_{\mu\nu}$ and of the matrices α_μ , associated with the wave equations of part I, is constructed. In this representation s , the total spin, is diagonal, which simplifies the calculation of the simultaneous mass and spin eigenvalues. Examples of mass-spin spectra are given, and it is proved that in certain, easily recognized, cases the mass eigenvalues are not all independent.

The matrix elements of the magnetic moment are calculated, and an example is given of a particle with an intrinsic magnetic moment equal to that of the proton.

1. INTRODUCTION

In part I (Le Couteur 1950, referred to as I) the discussion of linear relativistic wave equations of finite order was carried as far as possible without use of matrix representations of the α_μ and $I_{\mu\nu}$, which characterize the equations. This paper describes an explicit representation of these matrices.

It was shown in I (§ 6) that the physical properties of the wave equation appear most simply in a representation with α_0 and S , the total spin angular momentum, in diagonal form. Therefore, it is convenient to choose the basis of each representation (p, q) of the $I_{\mu\nu}$ so that S is diagonal as well as p and q ; then the coupling matrices which generate the α_μ can be worked out by the method indicated in § 8 of I. Finally, if desired, α_0 can be brought to diagonal form by a unitary transformation which commutes with S . This procedure differs from that of Dirac (1936) and Bhabha (1945) who expressed the coupling matrices in terms of spinor matrices $U^\alpha(k)$ and $V^\alpha(k)$. In that representation S is not diagonal and a complicated transformation, which has been worked out by Wild (1947), is required to diagonalize it.

The matrix representation of the α_μ directly relates the arbitrary coupling coefficients C_{ab} to the mass-spin spectrum and to the magnetic moment of the particle.

2. A REPRESENTATION OF THE LORENTZ GROUP

The representation with S diagonal has been worked out by Harish-Chandra (1947*a*). In the three-dimensional vector notation,

$$S_1 = iI_{23}, \text{ etc.}, \quad L_1 = I_{01}, \text{ etc.}, \quad (2.1)$$

the commutation rules (I, 2.8) for the $I_{\mu\nu}$ are

$$(S_1, S_2) = iS_3, \quad (2.2a)$$

$$(S_1, L_1) = 0, \quad (S_1, L_2) = iL_3 = -(S_2, L_1), \quad (2.2b)$$

$$(L_1, L_2) = iL_3. \quad (2.2c)$$

The well-known matrix representation of S is

$$\begin{aligned} (s, m | S_+ | s, m-1) &= \sqrt{\{(s-m+1)(s+m)\}}, \\ (s, m-1 | S_- | s, m) &= \sqrt{\{(s-m+1)(s+m)\}}, \\ (s, m | S_3 | s, m) &= m, \end{aligned} \quad (2.3)$$

where $S_{\pm} = S_1 \pm iS_2$ and $S^2 = s(s+1)$.

Any space-vector L which obeys the commutation rules (2.2*b*) can be expressed in terms of certain matrix elements given by Weyl (1931),

$$\left. \begin{aligned} (s, m | l_+ | s-1, m-1) &= -\sqrt{\{(s+m)(s+m-1)\}}, \\ (s, m-1 | l_- | s-1, m) &= \sqrt{\{(s-m)(s-m+1)\}}, \\ (s, m | l_3 | s-1, m) &= \sqrt{\{(s+m)(s-m)\}}, \end{aligned} \right\} \quad (2.4a)$$

$$\left. \begin{aligned} (s, m | l_+ | s, m-1) &= \sqrt{\{(s-m+1)(s+m)\}}, \\ (s, m-1 | l_- | s, m) &= \sqrt{\{(s-m+1)(s+m)\}}, \\ (s, m | l_3 | s, m) &= m, \end{aligned} \right\} \quad (2.4b)$$

$$\left. \begin{aligned} (s, m | l_+ | s+1, m-1) &= \sqrt{\{(s-m+1)(s-m+2)\}}, \\ (s, m-1 | l_- | s+1, m) &= -\sqrt{\{(s+m)(s+m+1)\}}, \\ (s, m | l_3 | s+1, m) &= \sqrt{\{(s+m+1)(s-m+1)\}}, \end{aligned} \right\} \quad (2.4c)$$

where $l_{\pm} = l_1 \pm il_2$.

Products of l matrices may be simplified by use of the following formulae:

$$\left. \begin{aligned} (s | (l_3, l_{\pm}) | s \pm 2) &= 0, \\ (s | (l_3, l_{\pm}) | s+1) \begin{Bmatrix} \text{via } s \\ \text{via } s+1 \end{Bmatrix} &= \pm \begin{Bmatrix} -s \\ s+2 \end{Bmatrix} (s | l_{\pm} | s+1), \\ (s | (l_3, l_{\pm}) | s-1) \begin{Bmatrix} \text{via } s-1 \\ \text{via } s \end{Bmatrix} &= \pm \begin{Bmatrix} 1-s \\ s+1 \end{Bmatrix} (s | l_{\pm} | s-1), \\ (s | (l_3, l_{\pm}) | s) \begin{Bmatrix} \text{via } s+1 \\ \text{via } s \\ \text{via } s-1 \end{Bmatrix} &= \pm \begin{Bmatrix} -(2s+3) \\ 1 \\ (2s-1) \end{Bmatrix} (s | l_{\pm} | s), \end{aligned} \right\} \quad (2.5)$$

where $(s | l_3 l_+ | s'')$ via s' means $(s | l_3 | s') (s' | l_\pm | s'')$ and so on. Also

$$\left. \begin{aligned} (s | \mathbf{1}^2 | s) \begin{Bmatrix} \text{via } s+1 \\ \text{via } s \\ \text{via } s-1 \end{Bmatrix} &= \begin{Bmatrix} (s+1)(2s+3) \\ s(s+1) \\ s(2s-1) \end{Bmatrix}, \\ (s | \mathbf{S1} | s) \begin{Bmatrix} \text{via } s+1 \\ \text{via } s \\ \text{via } s-1 \end{Bmatrix} &= \begin{Bmatrix} 0 \\ s(s+1) \\ 0 \end{Bmatrix}. \end{aligned} \right\} \quad (2.6)$$

In this notation the matrix elements of $I_{0k} = L_k$ can be expressed as

$$(s | L_k | s) = a(s) (s | l_k | s), \quad (s | L_k | s \pm 1) = a_\pm(s) (s | l_k | s \pm 1), \quad (2.7)$$

and the parameters $a(s)$ may be evaluated by a method due to Harish-Chandra (1947*a*). For the irreducible representation (p, q) of the $I_{\mu\nu}$, equations (I, 5.4) give

$$\left. \begin{aligned} \frac{1}{2}(\mathbf{S}^2 + \mathbf{L}^2) &= p(p+1) + q(q+1), \\ \mathbf{SL} &= p(p+1) - q(q+1). \end{aligned} \right\} \quad (2.8)$$

The left-hand side of (2.8) may be evaluated by use of (2.6) to give, with the notation,

$$\left. \begin{aligned} \lambda &= p+q, \quad \mu = p-q, \\ a(s) &= \frac{\mu(\lambda+1)}{s(s+1)}, \end{aligned} \right\} \quad (2.9)$$

$$a_+(s) = a_-(s+1) = \pm \sqrt{\left\{ \frac{[(s+1)^2 - \mu^2][(\lambda+1)^2 - (s+1)^2]}{(2s+1)(2s+3)(s+1)^2} \right\}}. \quad (2.10)$$

In this representation the matrices S_k and L_k are Hermitian. Since $S_3, L_3, S_\pm$ and $L_\pm$ are purely real, each of the $I_{\mu\nu}$ is either purely real or purely imaginary.

The sign factor in (2.10) cannot be determined by consideration of the commutation rules (2.2) only. However, in the representation $D(p, q)$ of the full Lorentz group, η_0 commutes with $\mathbf{S}$ and anti-commutes with $\mathbf{L}$. For $p \neq q$, according to (I, 5.6), the matrix elements of η_0 are chosen as $(p, q, s | \eta_0 | q, p, s) = 1$, and then in (2.10) we must take different signs for $p > q$ and $p < q$. For $p = q$ there is no such condition. Thus one can use the following convention:

$$\left. \begin{aligned} p \geq q, \text{ i.e. } \mu \geq 0 & \quad + \text{ sign in (2.10),} \\ p < q, \text{ i.e. } \mu < 0 & \quad - \text{ sign in (2.10).} \end{aligned} \right\} \quad (2.11)$$

In the (p, p) representation $\mu = 0$ and (2.10) shows that all the $a(s)$ vanish. In the representation $D(p, p)$ of the full Lorentz group $\eta_0^+(p, p)$ then has matrix elements

$$(s | \eta_0^+(p, p) | s) = (-)^{(s)}, \quad (2.12)$$

where $\{s\}$ denotes the integral part of s . It is easy to verify that this gives $\text{spur } \eta_0^+(p, p) = (-)^{2p} (2p+1)$, in agreement with (I, 5.19).

3. THE REPRESENTATION OF α_0

α_0 is to be determined from the conditions (I, 8.6)

$$(\alpha_0, I_{kl}) = 0 \quad ((\alpha_0, I_{0l}), I_{0l}) = \alpha_0. \quad (3.1)$$

The form of the coupling matrix $(a | \alpha_0 | b)$ can be worked out by considering only certain components of equation (3.1), but since, as was shown in § 8 of I, the coupling matrix exists and is unique, detailed consideration of the consistency of the other components is not necessary. Let I_{03}^{+-0} denote the parts of I_{03} which increase s by $+1, -1, 0$ respectively. Then, since α_0 commutes with s ,

$$((\alpha_0, I_{03}^0), I_{03}^+) + ((\alpha_0, I_{03}^+), I_{03}^0) = 0,$$

$$\text{or} \quad \alpha_0 [I_{03}^0, I_{03}^+] - 2I_{03}^0 \alpha_0 I_{03}^+ + [I_{03}^0, I_{03}^+] \alpha_0 - 2I_{03}^+ \alpha_0 I_{03}^0 = 0, \quad (3.2)$$

$$\text{where} \quad [x, y] = xy + yx. \quad (3.3)$$

Only matrix elements $(a | \alpha_0 | b)$ allowed by the condition (I, 8.1)

$$p_a - p_b = \pm \frac{1}{2}, \quad q_a - q_b = \pm \frac{1}{2}$$

need be considered.

(i) The coupling $(p, q) \leftrightarrow (p + \frac{1}{2}, q + \frac{1}{2})$ or $(\lambda, \mu) \leftrightarrow (\lambda + 1, \mu)$. All the terms in (3.2) involve a factor $m(s, m | l_3 | s + 1, m)$ which can be omitted, so that only the amplitude factors $a(s)$ of (2.7) need be taken into account. Equation (2.10) shows that a further factor $\{[(s+1)^2 - \mu^2]/(2s+1)(2s+3)(s+1)^2\}^{\frac{1}{2}}$ can be removed and so (3.2) gives

$$\begin{aligned} (\lambda, \mu, s | \alpha_0 | \lambda + 1, \mu, s) \{(\lambda + 2)^2 - (s + 1)^2\}^{\frac{1}{2}} \mu \left\{ \frac{\lambda + 2}{s(s+1)} + \frac{\lambda + 2}{(s+1)(s+2)} - \frac{2(\lambda + 1)}{s(s+1)} \right\} \\ + \{(\lambda + 1)^2 - (s + 1)^2\}^{\frac{1}{2}} \mu \left\{ \frac{\lambda + 1}{s(s+1)} + \frac{\lambda + 1}{(s+1)(s+2)} - \frac{2(\lambda + 2)}{(s+1)(s+2)} \right\} \\ \times (\lambda, \mu, s + 1 | \alpha_0 | \lambda + 1, \mu, s + 1) = 0, \end{aligned}$$

$$\text{or} \quad (\lambda, \mu, s | \alpha_0 | \lambda + 1, \mu, s) (\lambda + s + 3)^{\frac{1}{2}} (\lambda + 1 - s)^{\frac{1}{2}} (s - \lambda) \mu \\ + (\lambda, \mu, s + 1 | \alpha_0 | \lambda + 1, \mu, s + 1) (\lambda + s + 2)^{\frac{1}{2}} (\lambda - s)^{\frac{1}{2}} (\lambda + 1 - s) \mu = 0. \quad (3.4)$$

The other conditions contained in (3.1) must lead to equivalent equations. Some of these do not contain a common factor μ , so, even if $\mu = 0$, one can deduce

$$\frac{(\lambda, \mu, s | \alpha_0 | \lambda + 1, \mu, s)}{(\lambda + s + 2)^{\frac{1}{2}} (\lambda + 1 - s)^{\frac{1}{2}}} = \frac{(\lambda, \mu, s + 1 | \alpha_0 | \lambda + 1, \mu, s + 1)}{(\lambda + s + 3)^{\frac{1}{2}} (\lambda - s)^{\frac{1}{2}}}. \quad (3.5)$$

(ii) The coupling $(p, q) \leftrightarrow (p + \frac{1}{2}, q - \frac{1}{2})$ or $(\lambda, \mu) \leftrightarrow (\lambda, \mu + 1)$. The equations give

$$\frac{(\lambda, \mu, s | \alpha_0 | \lambda, \mu + 1, s)}{(\mu + s + 1)^{\frac{1}{2}} (s - \mu)^{\frac{1}{2}}} = \pm \frac{(\lambda, \mu, s + 1 | \alpha_0 | \lambda, \mu + 1, s + 1)}{(\mu + s + 2)^{\frac{1}{2}} (s + 1 - \mu)^{\frac{1}{2}}}. \quad (3.6)$$

The ambiguity of sign is a consequence of (2.11). The negative sign arises only when $\mu + 1 \geq 0$ and $\mu < 0$, that is for $\mu = -\frac{1}{2}$ and $\mu = -1$. Therefore, finally, one can take the matrix elements of α_0 in the form $(a | \alpha_0 | b) = C_{ab}(a | \alpha_0 | b)$ as

$$\left. \begin{aligned} (\lambda, \mu, s | \alpha_0 | \lambda + 1, \mu, s) &= (\lambda, \mu | C | \lambda + 1, \mu) (\lambda + s + 2)^{\frac{1}{2}} (\lambda + 1 - s)^{\frac{1}{2}}, \\ (\lambda + 1, \mu, s | \alpha_0 | \lambda, \mu, s) &= (\lambda + 1, \mu | C | \lambda, \mu) (\lambda + s + 2)^{\frac{1}{2}} (\lambda + 1 - s)^{\frac{1}{2}}, \\ (\lambda, \mu, s | \alpha_0 | \lambda, \mu + 1, s) &= (\lambda, \mu | C | \lambda, \mu + 1) (\mu + s + 1)^{\frac{1}{2}} (s - \mu)^{\frac{1}{2}} (-)^{(s)}, \\ (\lambda, \mu + 1, s | \alpha_0 | \lambda, \mu, s) &= (\lambda, \mu + 1 | C | \lambda, \mu) (\mu + s + 1)^{\frac{1}{2}} (s - \mu)^{\frac{1}{2}} (-)^{(s)}, \end{aligned} \right\} \quad (3.7)$$

and the sign factors $(-)^{(s)}$ appear only for $\mu = -\frac{1}{2}$ and $\mu = -1$.

These matrix elements agree with Wild's (1947) results apart from sign factors which can be removed by a change of basis. As was anticipated in § 8 of I, the coupling matrices are purely real and $(a|a_0|b)$ is the Hermitian conjugate of $(b|a_0|a)$.

The coupling coefficients $(a|C|b)$ or C_{ab} are not all independent, for the condition $\alpha_0\eta_0 = \eta_0\alpha_0$ gives

$$(\lambda, \mu|C|\lambda', \mu') = (\lambda, -\mu|C|\lambda', -\mu') \quad \text{for } \mu, \mu' \neq 0, \quad (3.8a)$$

and by (2.12),

$$(\lambda, 1|C|\lambda, 0) = \begin{cases} +(\lambda, -1|C|\lambda, 0) & \text{if } (\lambda, 0) \text{ refers to } \left\{ D_+\left(\frac{\lambda}{2}, \frac{\lambda}{2}\right) \right\} \\ -(\lambda, -1|C|\lambda, 0) & \text{the representation } \left\{ D_-\left(\frac{\lambda}{2}, \frac{\lambda}{2}\right) \right\} \end{cases} \quad (3.8b)$$

4. REALITY CONDITIONS

The physical interpretation requires (I, § 6) that α_0 can be reduced to diagonal form with real coefficients. Thus α_0 and, equally, all the α_μ , must be equivalent to its Hermitian conjugate $\alpha_0^\dagger$, conjugate imaginary $\bar{\alpha}_0$ and transposed $\tilde{\alpha}_0$. The first requirement is already embodied in the equation (I, 3.2),

$$\alpha_\mu^\dagger D = D\alpha_\mu, \quad (4.1)$$

which gives a condition (I, 8.8)

$$\bar{C}_{ba} = K_a C_{ab} K_b \quad (4.2)$$

on the coupling coefficients. The others are satisfied if there exists a matrix F such that

$$\bar{\alpha}_\mu = F\alpha_\mu F^{-1} \quad \text{or} \quad \bar{\alpha}_\mu F = F\alpha_\mu, \quad (4.3)$$

but this condition is not strictly necessary, since one could have a different F for each α_μ . However, (4.3) is a useful restriction on the mathematical possibilities. Now (Harish-Chandra 1947*b*) the matrices $I_{\rho\sigma}$ and η_0 are uniquely determined by the α_μ through the commutation rules (I, 2.6) and (I, 2.9); similarly, $\bar{I}_{\rho\sigma}$ and $\bar{\eta}_0$ are determined by the $\bar{\alpha}_\mu$. Therefore, (4.2) requires

$$\bar{I}_{\rho\sigma} = F I_{\rho\sigma} F^{-1}, \quad \bar{\eta}_0 = F \eta_0 F^{-1}. \quad (4.4)$$

The commutation rules (I, 2.8), which define the structure of the Lorentz group, show that a matrix F satisfying (4.4) must exist independently of the assumption made in (4.3).

Equations (4.3) and (4.4) show that $\bar{F}F$ and $F\bar{F}$ commute with all the α_μ and $I_{\mu\nu}$ and so, in an irreducible system of equations, they must be equal, real multiples of the unit matrix. Only the sign is determinate and so one can always have

$$F\bar{F} = \bar{F}F = +1 \text{ or } -1 \quad (4.5)$$

The definitions (2.8) of p and q show that the matrix elements of F are of the form $(p, q, s|F|q, p, s)$. Since in the representation of § 2, S_3 , L_3 , $S_\pm$, $L_\pm$ and η_0 are purely real, F commutes with I_{31} , I_{03} , I_{01} , η_0 and anticommutes with I_{12} , I_{23} , I_{02} .

These conditions fully determine the form of the submatrix of F contained in the representation $D(p, q)$ as

$$(p, q, s, m | F | q, p, s, -m) = f(p, q) \begin{cases} (-)^{s+m} & \text{for } p \neq q \\ (-)^m & \text{for } p = q \end{cases}, \quad (4.6)$$

where $f(p, q) = f(q, p)$ and (4.5) gives $|f(p, q)| = 1$. In this representation F closely resembles the improper transformation η_2 which reverses the x_2 axis.

These results can be used to prove the theorem that, if no representation (p, q) occurs more than once with the wave equation, the coupling coefficients C_{ab} or $(\lambda, \mu | C | \lambda', \mu')$ can all be transformed into purely real numbers. This simplifies practical calculations and also shows that in such cases each coupling coefficient represents only a single degree of freedom.

The result follows from (4.3) which, applied to the representation (3.7) of α_0 , gives

$$(\lambda, \mu | C | \lambda', \mu') f(\lambda', \mu') = f(\lambda, \mu) (\lambda, -\mu | C | \lambda', -\mu'). \quad (4.7)$$

Let
$$f(\lambda, \mu) = \pm e^{i\theta(\lambda, \mu)}, \quad (4.8)$$

where the negative sign occurs only if $\mu = 0$ and $(\lambda, 0)$ belongs to the representation $D-(\frac{1}{2}\lambda, \frac{1}{2}\lambda)$. Then, by use of (3.8), (4.7) becomes

$$e^{-\frac{1}{2}i\theta(\lambda, \mu)} (\lambda, \mu | C | \lambda', \mu') e^{\frac{1}{2}i\theta(\lambda', \mu')} = e^{\frac{1}{2}i\theta(\lambda, \mu)} (\lambda, \mu | C | \lambda', \mu') e^{-\frac{1}{2}i\theta(\lambda', \mu')}, \quad (4.9)$$

which shows that the similarity transformation

$$\alpha_\mu \rightarrow e^{\frac{1}{2}i\theta} \alpha_\mu e^{-\frac{1}{2}i\theta} \quad (4.10)$$

reduces all the coupling coefficients C_{ab} to purely real form. Since $f(p, q) = f(q, p)$ the representation of η_0 and of D is unchanged by this transformation.

For a large class of wave equations the existence of a matrix F satisfying (4.3) can be proved. Suppose that the coupling diagram (I, § 10) contains no closed loops lying wholly to one side or the other of the axis of symmetry $p = q$. Then the half-diagram $p \geq q$ consists of an open chain of representations (p_a, q_a) like

$$\begin{array}{c} a \text{---} b \text{---} d \\ | \\ c \end{array} \quad (4.11)$$

and by a transformation
$$\alpha_\mu \rightarrow e^{i\phi} \alpha_\mu e^{-i\phi} \quad (4.12)$$

with
$$\phi_a = 0, \quad \phi_b = \arg C_{ab}, \quad \phi_d = \phi_b + \arg C_{bd} \text{ and so on,} \quad (4.13)$$

all the coupling coefficients C_{ab} in this half-diagram become real numbers. If, for the representations in the other half-diagram, $p \leq q$, ϕ is defined by $\phi(p, q) = \phi(q, p)$, the transformation (4.12) leaves η_0 unchanged and then (3.8) shows that the coupling coefficients in this half-diagram also are real numbers. Having transformed all the coupling coefficients to real form, one can define the matrix F through (4.6), (4.7), (4.8) with $\theta(\lambda, \mu) = 0$. This F satisfies $\bar{\alpha}_0 F = F \alpha_0$ as well as (4.4), and these together imply (4.3).

If the half-diagram contains closed loops, the definition (4.13) of ϕ is not self-consistent except for special values of the C_{ab} . Thus, in this more general case, (4.3) must be regarded as an additional assumption.

5. THE REPRESENTATION OF (α_μ, α_ν)

In all simple theories the antisymmetric tensor (α_μ, α_ν) is closely related to, if not actually a multiple of, $I_{\mu\nu}$. The matrix elements of these brackets are useful. It follows from I, § 4, that the matrix elements $(a | (\alpha_\mu, \alpha_\nu) | b)$ connecting representations a and b are non-zero only if the representations $(1, 0)$ or $(0, 1)$ appear in the reduction (I, 5.13) of the direct product $(p_a, q_a) \times (p_b, q_b)$. Thus the possibilities are

$$p_a = p_b, \quad q_a = q_b; \quad (5.1)$$

$$\text{or} \quad p_a = p_b, \quad q_a = q_b \pm 1 \quad \text{or} \quad p_a = p_b \pm 1, \quad q_a = q_b. \quad (5.2)$$

The first possibility is of special interest for, since the direct product

$$(p_a, q_a) \times (p_a, q_a) = (0, 0) + (0, 1) + (1, 0) + \dots,$$

contains $(1, 0)$ and $(0, 1)$ once only, $(a | i(\alpha_2, \alpha_3) \pm (\alpha_0, \alpha_1) | a)$ must be a multiple $(x \pm y)$ or $(a | iI_{23} \pm I_{01} | a)$. Thus

$$\left. \begin{aligned} (a | (\alpha_0, \alpha_1) | a) &= x(a | I_{01} | a) + iy(a | I_{23} | a), \\ i(a | (\alpha_2, \alpha_3) | a) &= ix(a | I_{23} | a) + y(a | I_{01} | a). \end{aligned} \right\} \quad (5.3)$$

The coefficients x and y are easily calculated by the method of § 3, as

$$\begin{aligned} \frac{1}{2}x &= - |(\lambda, \mu | C | \lambda + 1, \mu)|^2 (\lambda + 2) + |(\lambda, \mu | C | \lambda - 1, \mu)|^2 \lambda \\ &\quad + |(\lambda, \mu | C | \lambda, \mu + 1)|^2 (\mu + 1) - |(\lambda, \mu | C | \lambda, \mu - 1)|^2 (\mu - 1), \end{aligned} \quad (5.4)$$

$$\begin{aligned} \frac{1}{2}y &= |(\lambda, \mu | C | \lambda + 1, \mu)|^2 \mu - |(\lambda, \mu | C | \lambda - 1, \mu)|^2 \mu \\ &\quad - |(\lambda, \mu | C | \lambda, \mu + 1)|^2 (\lambda + 1) + |(\lambda, \mu | C | \lambda, \mu - 1)|^2 (\lambda + 1), \end{aligned} \quad (5.5)$$

and the factors $|a | c | b|^2$ appear because products like

$$(\lambda, \mu | C | \lambda + 1, \mu) (\lambda + 1, \mu | C | \lambda, \mu)$$

have been simplified by use of (4.2) with $K = 1$ for simplicity. The matrix elements of (α_μ, α_ν) which couple two inequivalent representations of the Lorentz group cannot be expressed in terms of anything more elementary. However, the dependence on the coupling coefficients is simple. $(\lambda, \mu | (\alpha_\mu, \alpha_\nu) | \lambda + 1, \mu + 1)$ is proportional to

$$\begin{aligned} &(\lambda, \mu | C | \lambda + 1, \mu) (\lambda + 1, \mu | C | \lambda + 1, \mu + 1) (\lambda - \mu + 2) \\ &\quad + (\lambda, \mu | C | \lambda, \mu + 1) (\lambda, \mu + 1 | C | \lambda + 1, \mu + 1) (\mu - \lambda) \end{aligned} \quad (5.6)$$

and $(\lambda, \mu | (\alpha_\mu, \alpha_\nu) | \lambda + 1, \mu - 1)$ is proportional to

$$\begin{aligned} &(\lambda, \mu | C | \lambda + 1, \mu) (\lambda + 1, \mu | C | \lambda + 1, \mu - 1) (-\mu - \lambda - 2) \\ &\quad + (\lambda, \mu | C | \lambda, \mu - 1) (\lambda, \mu - 1 | C | \lambda + 1, \mu - 1) (\mu + \lambda). \end{aligned} \quad (5.7)$$

These relations can, for example, be used to determine the coupling coefficients in Bhabha's (1945) theory in which $I_{\mu\nu} = (\alpha_\mu, \alpha_\nu)$. More generally, if the coupling coefficients are assigned, the equations can be used to work out the connexion between the α_μ and the $I_{\mu\nu}$, and then the fundamental commutation rules for the $I_{\mu\nu}$ become commutation rules for the α algebra.

6. MASS-SPIN EIGENVALUES

The results of §§ 3 and 4 lead to some general conclusions about the relationship of the mass-spin eigenvalues to the coupling coefficients.

(i) The states of highest spin $s = S_H$. For simplicity a subdiagram of the form (I, 10.2),

$$(p, q) \leftrightarrow \cdots \leftrightarrow (q, p) \quad (6.1)$$

is considered. This chain of representations involves $2(p - q)$ coupling coefficients which, according to § 4, may be assumed to be real. Because of the reflexion symmetry, the number of independent parameters is $p - q$ or $p - q + \frac{1}{2}$ according as the spin is integral or half-odd integral. In either case, the submatrix $(S_H | \alpha_0 | S_H)$ has the form

$$\begin{pmatrix} 0 & a & & & \\ \tilde{a} & 0 & b & & \\ & \tilde{b} & 0 & & \\ & & & \ddots & \\ & & & & a \\ & & & & \tilde{a} & 0 \end{pmatrix} \quad (6.2)$$

the rows and columns being labelled by the representations in (6.1) and, according to (4.2), $\tilde{a} = \pm a$. The $2(p - q) + 1$ real eigenvalues of this submatrix of α_0 are

(a) for integral spins: $0, \pm j_r$,

(b) for half-odd spins: $\pm j_r$,

and so, in both cases, the number of independent coupling parameters is equal to the number of independent eigenvalues. Thus, using the physical interpretation of I, § 6, one can say that the coefficients which couple the representations (p, q) containing spin S_H are determined by the set of mass values associated with this, the highest, spin eigenvalue.

(ii) Consider a wave equation which involves only representations with $p + q = \text{constant} = S$. The above argument shows that the set of masses associated with spin S determines the whole set of coupling coefficients and therefore the whole mass-spin spectrum.

Example 1: integral spin. Consider the chain

$$(S + \tfrac{1}{2}, S - \tfrac{1}{2}) \leftrightarrow (S, S) \leftrightarrow (S - \tfrac{1}{2}, S + \tfrac{1}{2})$$

with coupling coefficients a .

Equations (3.7) give

$$(s | \alpha_0 | s) = a \begin{pmatrix} 0 & \sqrt{s(s+1)} & 0 \\ \sqrt{s(s+1)} & 0 & (-)^s \sqrt{s(s+1)} \\ 0 & (-)^s \sqrt{s(s+1)} & 0 \end{pmatrix}, \quad (6.3)$$

and so the mass associated with the spin eigenvalue $s \neq 0$ is

$$\chi/a\sqrt{\{2s(s+1)\}}.$$

Example 2: half-odd integral spin. Consider the chain

$$(S + \tfrac{1}{2}, S) \leftrightarrow (S, S + \tfrac{1}{2}).$$

Equations (3.7) give

$$(s | \alpha_0 | s) = \begin{pmatrix} 0 & (-)^s(s + \frac{1}{2}) \\ (-)^s(s + \frac{1}{2}) & 0 \end{pmatrix}, \quad (6.4)$$

and so the mass associated with spin s is $\chi/(s + \frac{1}{2})$.

(iii) The states of spin $S_H - 1$. In the general case one must consider the representations in the line $p + q = S_H - 1$, as well as those, already considered in (6.1), in the line $p + q = S_H$. Let us suppose that there is a complete chain

$$(p', q') \leftrightarrow \dots \leftrightarrow (q', p'), \quad p' + q' = S_H - 1, \quad (6.5)$$

of representations with $p + q = S_H - 1$. It contains $2(p' - q')$ coupling coefficients. There are also $2(p' - q') + 1$ or $2(p - q) + 1$, whichever is less, possible couplings between this line and the line (6.1) of representations containing spin S_H . Altogether the spin representation $D(S_H - 1)$ occurs $\{2(p' - q') + 1 + 2(p - q) + 1\}$ times and there are

$$p - q + p' - q' + \min\{(p' - q'), (p - q)\} + \begin{cases} 1 & \text{for } s \text{ integral} \\ 1\frac{1}{2} & \text{for } s \text{ half-odd} \end{cases}$$

independent coupling coefficients, of which the first $p - q$ or $p - q + \frac{1}{2}$ have been determined from the masses associated with spin S_H . Therefore, if $p' - q' < p - q$ the masses with spin $S_H - 1$ cannot be chosen quite independently of those with spin S_H . If $p' - q' > p - q$ the masses with spin $S_H - 1$ are independent of those with spin S_H and suffice to determine the additional coupling coefficients.

Only simple forms of the S_H and $S_H - 1$ subdiagrams have been considered. In other cases, slight modifications of these results are required.

(iv) One can continue in this way, working through the coupling diagram from the states of highest towards the states of lower spin. The general conclusion is:

If the total number of representations containing D_s is greater than the total number of couplings connecting the line $p + q = s$ to itself and to the line $s + 1$, then the masses associated with spin s cannot be chosen independently of the masses associated with states of higher spin. If the number of representations is odd it must be decreased by one and if the number of couplings is odd it must be increased by one, before applying this rule.

A general method has now been established which allows one to determine the coupling coefficients if the mass-spin spectrum is known. Of course, this is made possible by the assumption of § 9 of I that α_0 and D_s form a 'complete set' of variables, and the consequent restriction (§ 10) on the form of the coupling diagrams.

It is not possible, by this method, to work upwards from the states of spin 0 or $\frac{1}{2}$ towards the states of higher spin. For example, the representations containing spin 0 form a chain $(p, p) - (p - \frac{1}{2}, p - \frac{1}{2}) \dots (q, q)$ involving $2(p - q)$ couplings which are independent since they are not connected by reflexion symmetry. The number of masses associated with spin 0 is, however, only $p - q$ or $p - q + \frac{1}{2}$, and so they cannot be used to determine the couplings. Similar considerations apply if the lowest spin value is $\frac{1}{2}$.

(v) An example in which the mass eigenvalues are all independent and determine the coupling coefficients is the (p, q) diagram

$$\begin{array}{ccc} (1, \frac{1}{2}) & \text{---} c \text{---} & (\frac{1}{2}, 1) \\ | & & | \\ b & & b \\ | & & | \\ (\frac{1}{2}, 0) & \text{---} a \text{---} & (0, \frac{1}{2}) \end{array} \quad (6.6)$$

The matrix elements (3.7) give for α_0 ,

$$(s = \frac{3}{2} | \alpha_0 | s = \frac{3}{2}) = \begin{pmatrix} 0 & -2c \\ -2c & 0 \end{pmatrix},$$

$$(s = \frac{1}{2} | \alpha_0 | s = \frac{1}{2}) = \begin{pmatrix} 0 & b\sqrt{3} & 0 & c \\ b\sqrt{3} & 0 & a & 0 \\ 0 & a & 0 & b\sqrt{3} \\ c & 0 & b\sqrt{3} & 0 \end{pmatrix}$$

and the eigenvalues j of α_0 are therefore

$$\left. \begin{array}{l} \text{for } s = \frac{3}{2} \quad j = \pm 2c, \\ \text{for } s = \frac{1}{2} \quad j = \pm \frac{1}{2} \{a + c \pm [(a - c)^2 + 12b^2]^{\frac{1}{2}}\} = \pm v \text{ or } w, \end{array} \right\} \quad (6.7)$$

and, as usual, the mass eigenvalues are $|\chi/j|$. The scheme (6.6) has the same structure as Bhabha's (1945) equation $R_5(\frac{3}{2}, \frac{1}{2})$, in which the eigenvalues are

$$\text{for } s = \frac{3}{2}, \quad j = \pm \frac{1}{2}; \quad \text{for } s = \frac{1}{2}, \quad j = \pm \frac{1}{2}, \quad \pm \frac{3}{2}.$$

By substituting these values into (6.7) the coupling coefficients are determined as

$$\left. \begin{array}{l} c = \frac{1}{4}, \quad a = \frac{3}{4}, \quad b = \frac{1}{4}\sqrt{5} \\ c = \frac{1}{4}, \quad a = -\frac{5}{4}, \quad b = \frac{1}{4}\sqrt{3}. \end{array} \right\} \quad (6.8)$$

The first solution agrees with the values determined by Bhabha from his commutation rules, allowance being made for the fact that his coupling coefficients are larger, by a factor 2, than those defined by (3.7).

It is easy to choose the constants a, b, c in the diagram (6.6) so that, in the stable state of lowest rest mass, the particle has spin $\frac{1}{2}$. Then this equation might perhaps be used to represent a proton or neutron, but it would imply that these particles possess states of higher mass, with spin $\frac{3}{2}$ and $\frac{1}{2}$.

7. THE MAGNETIC MOMENT

It is now possible to calculate the matrix elements of the magnetic moment (I, 11.9),

$$\frac{e}{M} \frac{i}{2j} (j | I_{02}\alpha_3 - I_{03}\alpha_2 | j). \quad (7.1)$$

Now

$$I_{02}\alpha_3 - I_{03}\alpha_2 = I_{02}\alpha_0 I_{03} - I_{03}\alpha_0 I_{02} + I_{23}\alpha_0. \quad (7.2)$$

Then, by use of (2.5),

$$\begin{aligned}
 & (\lambda, \mu, s | I_{02} \alpha_0 I_{03} - I_{03} \alpha_0 I_{02} | \lambda + 1, \mu, s) \\
 &= a(\lambda, \mu, s) a(\lambda + 1, \mu, s) (\lambda, \mu, s | \alpha_0 | \lambda + 1, \mu, s) \\
 & \quad + a_+(\lambda, \mu, s) a_-(\lambda + 1, \mu, s + 1) (\lambda, \mu, s + 1 | \alpha_0 | \lambda + 1, \mu, s + 1) \{-(2s + 3)\} \\
 & \quad + a_-(\lambda, \mu, s) a_+(\lambda + 1, \mu, s - 1) (\lambda, \mu, s - 1 | \alpha_0 | \lambda + 1, \mu, s - 1) (2s - 1) \Big\} (s | -I_{23} | s) \\
 &= (\lambda, \mu | C | \lambda + 1, \mu) \{(\lambda + s + 2)(\lambda + 1 - s)\}^{\frac{1}{2}} \left\{ 2 - \frac{\mu^2}{s(s+1)} \right\} (s | -I_{23} | s), \quad (7.3)
 \end{aligned}$$

by use of (3.7) and (2.9). Similarly,

$$\begin{aligned}
 & (\lambda, \mu, s | I_{02} \alpha_0 I_{03} - I_{03} \alpha_0 I_{02} | \lambda, \mu + 1, s) \\
 &= (\lambda, \mu | C | \lambda, \mu + 1) \{(\mu + s + 1)(s - \mu)\}^{\frac{1}{2}} \left\{ 2 - \frac{(\lambda + 1)^2}{s(s+1)} \right\} (s | -I_{23} | s). \quad (7.4)
 \end{aligned}$$

Therefore, by (7.1), the required result is

$$\begin{aligned}
 & (\lambda, \mu, s | I_{02} \alpha_3 - I_{03} \alpha_2 | \lambda + 1, \mu, s) \\
 &= (\lambda, \mu | C | \lambda + 1, \mu) \left\{ \frac{\mu^2}{s(s+1)} - 1 \right\} \{(\lambda + s + 2)(\lambda + 1 - s)\}^{\frac{1}{2}} (s | I_{23} | s), \quad (7.5)
 \end{aligned}$$

and $(\lambda, \mu, s | I_{02} \alpha_3 - I_{03} \alpha_2 | \lambda, \mu + 1, s)$

$$= (\lambda, \mu | C | \lambda, \mu + 1) \left\{ \frac{(\lambda + 1)^2}{s(s+1)} - 1 \right\} \{(\mu + s + 1)(s - \mu)\}^{\frac{1}{2}} (s | I_{23} | s).$$

There are also matrix elements of the form $(\lambda, \mu, s | I_{02} \alpha_3 - I_{03} \alpha_2 | \lambda', \mu', s \pm 1)$, but, to determine the magnetic moment in the state of lowest rest mass, it is not necessary to evaluate these explicitly.

Example. It is of some interest to work out the magnetic moment for the wave equation discussed in § 6 (v). It is easy to choose the coupling coefficients so as to give the observed magnetic moment, which, as is well known, is not correctly given by Dirac's equation.

The matrix elements (7.5) give

$$\frac{1}{2}(s = \frac{1}{2} | I_{02} \alpha_3 - I_{03} \alpha_2 | s = \frac{1}{2}) = \begin{pmatrix} 0 & -b/\sqrt{3} & 0 & 11c/3 \\ -b/\sqrt{3} & 0 & a & 0 \\ 0 & a & 0 & -b/\sqrt{3} \\ 11c/3 & 0 & -b/\sqrt{3} & 0 \end{pmatrix} (\frac{1}{2} | I_{23} | \frac{1}{2}), \quad (7.6)$$

and the relevant part of α_0 is

$$(s = \frac{1}{2} | \alpha_0 | s = \frac{1}{2}) = \begin{pmatrix} 0 & b\sqrt{3} & 0 & c \\ b\sqrt{3} & 0 & a & 0 \\ 0 & a & 0 & b\sqrt{3} \\ c & 0 & b\sqrt{3} & 0 \end{pmatrix} \begin{matrix} 1 \\ 2 \\ 3 \\ 4 \end{matrix}. \quad (7.7)$$

For convenience the rows and columns of these matrices are labelled 1, 2, 3, 4, which correspond to the representations $(1, \frac{1}{2})$, $(\frac{1}{2}, 0)$, $(0, \frac{1}{2})$, $(\frac{1}{2}, 1)$ contained in the diagram (6.6). The eigenvectors of α_0 of the form $\psi_1 = \psi_4$, $\psi_2 = \psi_3$ are given by

$$(j - c) \psi_1 - b\sqrt{3} \psi_2 = 0, \quad -b\sqrt{3} \psi_1 + (j - a) \psi_2 = 0, \quad (7.8)$$

with

$$j = \frac{1}{2}(c+a) \pm \{(c-a)^2 + 12b^2\}^{\frac{1}{2}},$$

and the eigenvectors of the form $\psi_1 = -\psi_4$, $\psi_2 = -\psi_3$ lead to the opposite pair of eigenvalues. The component of (7.6) corresponding to the eigenvalue j of α_0 is given by

$$\frac{\begin{pmatrix} b\sqrt{3} \\ j-c \\ j-c \\ b\sqrt{3} \end{pmatrix} \begin{pmatrix} 0 & -b/\sqrt{3} & 0 & 11c/3 \\ -b/\sqrt{3} & 0 & a & 0 \\ 0 & a & 0 & -b/\sqrt{3} \\ 11c/3 & 0 & -b/\sqrt{3} & 0 \end{pmatrix} \begin{pmatrix} b\sqrt{3} \\ j-c \\ j-c \\ b\sqrt{3} \end{pmatrix}}{(6b^2 + 2(j-c)^2)} = \frac{-2b^2(j-c) + 11b^2c + a(j-c)^2}{3b^2 + (j-c)^2} = k_j. \quad (7.9)$$

Therefore, according to (7.1), the magnetic moment in the state of mass $M = \chi/j$ and spin $\frac{1}{2}$ is given by

$$k \frac{e}{M} (S_1, S_2, S_3) \quad (7.10)$$

with the anomalous factor k .

Let v be the larger, and w the smaller eigenvalue of j given by (7.8). Then, for the state of lowest rest mass, with $j = v$, equation (7.9) gives

$$k = 1 + \frac{8(v-2c)(w-c)}{3v(v-w)}, \quad (7.11)$$

and the mass spectrum of the particle is χ/v , χ/w , $\chi/2c$. Thus (7.12) expresses a direct relationship between the mass spectrum and the magnetic moment.

As a check, note that if $b = c = 0$ the equation reduces to Dirac's equation; then (7.8) gives $w = 0$ and $k = 1$, in agreement with the well-known result.

8. WAVE EQUATIONS FOR THE NUCLEONS

The proton and neutron have magnetic moment $M_p = 2.78$ and $M_n = -1.9$ proton magnetons. It is thought that these anomalous values arise because of the interaction between the nucleons and the meson field, but no satisfactory theory is available.

As an application of the formulae of § 7, let us determine a wave equation which yields an intrinsic magnetic moment equal to that of the proton. In (7.11), this requires $k = 2.78$ or

$$\frac{(v-2c)(w-c)}{v(v-w)} = \frac{3}{8} \times 1.78 = 0.66. \quad (8.1)$$

Only the ratios of v , w , $2c$ matter. It is convenient to write $v = 1$, $2c = 1 - \epsilon$ and then, for small ϵ , equation (8.1) gives $w = 1 - 0.75\epsilon$. Thus, in this approximation, the mass spin spectrum (6.7) becomes

$$\left. \begin{array}{lll} \text{mass } 1 & 1 + 0.75\epsilon & 1 + \epsilon \\ \text{spin } \frac{1}{2} & \frac{1}{2} & \frac{3}{2} \end{array} \right\}. \quad (8.2)$$

If ϵ is positive, at low energies the particle behaves as if it had only a single spin of magnitude $\frac{1}{2}$ and in addition it shows the correct intrinsic magnetic moment. The higher mass and spin states included in (8.2) are an essential feature of the model but there is not yet sufficient experimental evidence about the behaviour of high energy protons to say whether such excited states exist.

However, according to present ideas, it is not reasonable to discuss the magnetic moment of the proton without at the same time giving an account of that of the neutron. With meson theory, the ideal procedure for these problems which do not involve free mesons would be to start from the Hamiltonian for a system of nucleons, mesons and photons in interaction and to eliminate the interaction with the mesons by canonical transformation. The improved methods of Schwinger (1949) and Dyson (1949) can be applied to such problems and must lead to a modified Hamiltonian for the system of nucleons and photons only. The one nucleon part will reduce to the form

$$H_p + H'_p, \quad \text{or} \quad H_n + H'_n, \quad (8.3)$$

where H_p and H_n are the effective Hamiltonians for a free proton or neutron and H'_p and H'_n additional Hamiltonians for interaction with the electromagnetic field.

The transformed Hamiltonians H_n and H_p must satisfy the requirements of relativistic invariance. However, the elimination of the meson interaction is likely to involve the appearance of integral operators, or of high derivatives of the nucleon wave function. If the resulting wave equation can be approximated to by a differential equation of finite order, it can presumably be thrown into the form

$$(\alpha^\mu p_\mu - \chi) \psi = 0,$$

discussed in these papers. But the above considerations show that the resulting α^μ may well be more complicated than the Dirac matrices; so that the formal theory of generalized wave equations is relevant to this problem. The strong-coupling theory (see Pauli 1946) provides additional support for this contention, for the theory predicts that the strong meson-nucleon interaction must lead to the existence of excited states of the nucleons, with higher spin and effective mass. As has already been discussed (I, § 6), the most typical feature of all generalized wave equations is that they describe particles having just such a spectrum of mass and spin values.

The magnetic moment of the nucleons is made up of two parts, first an explicit contribution from the transformed interaction Hamiltonians H'_n and H'_p and secondly for the proton there is the intrinsic magnetic moment which can be calculated from H_p as in § 7. The first contribution has been calculated by many authors, most recently by Case (1949). According to him it is difficult to get simultaneous quantitative agreement with the experimental moments of proton and neutron by using this first term together with an intrinsic moment of 1 proton magneton for the proton. This may well be owing to the inadequacy of present meson theory: still it shows that there is room for an 'anomalous' value of the kind discussed in § 7, for the intrinsic magnetic moment of the proton.

If these, theoretically plausible, excited states of the nucleons do in fact exist, they should be observable, for example, by scattering experiments at high energies.

When the states are known, it becomes a fairly straightforward matter to derive the appropriate wave equation by the methods of this paper, and then the intrinsic magnetic moment can also be calculated by the methods described.

REFERENCES

- Bhabha, H. J. 1945 *Rev. Mod. Phys.* **17**, 200.
 Case, K. M. 1949 *Phys. Rev.* **76**, 1.
 Dirac, P. A. M. 1936 *Proc. Roy. Soc. A*, **155**, 447.
 Dyson, F. J. 1949 *Phys. Rev.* **75**, 486.
 Harish-Chandra 1947*a* *Proc. Roy. Soc. A*, **189**, 372.
 Harish-Chandra 1947*b* *Phys. Rev.* **71**, 793.
 Le Couteur, K. J. 1950 *Proc. Roy. Soc. A*, **202**, 284.
 Pauli, W. 1946 *Meson theory of nuclear forces*. New York: Interscience.
 Schwinger, J. 1949 *Phys. Rev.* **74**, 1439; **75**, 651.
 Weyl H. 1931 *Group theory and quantum mechanics*. London: Methuen.
 Wild, E. 1947 *Proc. Roy. Soc. A*, **191**, 253.

Finite elastic deformation of incompressible isotropic bodies

BY A. E. GREEN and R. T. SHIELD

Mathematics Department, King's College, Newcastle-upon-Tyne

(Communicated by G. R. Goldsbrough, F.R.S.—Received 22 February 1950)

The theory of finite elastic deformation of incompressible isotropic bodies is expressed in a simple form in tensor notation, using a general system of co-ordinates which move with the body as it is deformed. In order to illustrate the advantages of the present methods one problem, previously solved by Rivlin, is re-examined. Two new problems are then solved using a completely general form for the strain energy function, the first problem being that of a rotating cylinder, and the second a uniform spherical shell under symmetrical internal and external pressures.

1. INTRODUCTION

In a recent series of papers Rivlin (1948*a, b, c, d*, 1949*a, b, c*) has discussed at length the theory of finite deformation of isotropic elastic solids, with special reference to incompressible materials. Equations of motion, boundary conditions, and relations between components of stress and a strain-energy function, were obtained in a variety of equivalent forms, some of which had been given earlier by other writers. Rivlin then considered certain forms for the strain-energy function for incompressible solids, including one for a neo-Hookean solid which reduces in the case of infinitesimal strains and displacements to a classical form, and he solved a number of special problems.

The general theory is first developed by Rivlin in a special set of rectangular Cartesian co-ordinates (x, y, z) in the undeformed body. Relations between stress and the strain-energy function are expressed in terms of nine partial derivatives of

the displacement components (u, v, w) with respect to (x, y, z) , and the relations are later transformed and given in terms of strain invariants. Equations are also transformed to cylindrical polar co-ordinates. While deriving considerable help from Rivlin's work we feel that the whole theory is unnecessarily elaborate and complicated. This is partly because the theory is developed in terms of partial derivatives of displacements instead of in symmetrical strain components, and partly because of the choice of co-ordinates and the avoidance of any kind of shorthand notation.

In the present paper we reconsider the theory of finite deformation of isotropic incompressible elastic bodies. A general system of co-ordinates, and tensor notation, are used, and we believe that the theory then takes a simple and elegant form in which complicated algebraic manipulations are avoided without losing sight of the essential physical features in the theory. In a recent paper* Green & Zerna (1950) have presented an account of elasticity theory for compressible elastic materials in tensor and vector notations, using a system of moving co-ordinates, and the work for incompressible materials is a natural sequel to this. Tensor notations have, of course, been used previously by other writers (see, for example, Murnaghan (1937)), but the notation and methods used here appear to be rather simpler.

In order to illustrate the present form of the theory a problem solved by Rivlin is briefly examined, and other new problems are solved, using completely general forms for the strain-energy function.

2. NOTATION

We use the notations developed in A and briefly repeat the essential formulae here.

The points P_0 of an unstrained and unstressed elastic body at rest, called briefly the body B_0 , are defined at the time $t = 0$ by a rectangular Cartesian system of co-ordinates x_i or by a general curvilinear system of co-ordinates θ_i , where

$$x_i = x_i(\theta_1, \theta_2, \theta_3), \quad (2.1)$$

and the line element ds_0 of B_0 is given by

$$ds_0^2 = dx^i dx^i = g_{ik} d\theta^i d\theta^k. \quad (2.2)$$

The usual summation convention is used and Latin indices take the values 1, 2, 3. If an index is repeated more than twice it is not summed.

The body B_0 is now strained or deformed so that at time t the points P_0 of B_0 have moved to new positions P , and form a strained body called briefly the body B . The points P of B are defined by a new rectangular Cartesian system of co-ordinates y_i which are related to x_i by the equations

$$y_i = f_i(x_1, x_2, x_3, t). \quad (2.3)$$

The curvilinear co-ordinates θ_i in B_0 form a curvilinear system θ_i in B so that

$$y_i = y_i(\theta_1, \theta_2, \theta_3, t), \quad (2.4)$$

and the line element ds in B , for a given time t , is given by

$$ds^2 = dy^i dy^i = G_{ik} d\theta^i d\theta^k. \quad (2.5)$$

* We refer to this paper as A.

In the curvilinear system θ_i covariant and contravariant base vectors $\mathbf{e}_i, \mathbf{e}^i$ can be defined at each point P_0 of B_0 , and covariant and contravariant base vectors $\mathbf{E}_i, \mathbf{E}^i$ can be defined at each point P of B .

We define covariant components of a strain tensor as

$$\gamma_{ik} = \frac{1}{2}(G_{ik} - g_{ik}). \quad (2.6)$$

The mixed components can be defined in two ways, but for our purpose we choose

$$\gamma_s^r = g^{rl}\gamma_{sl}. \quad (2.7)$$

It was shown in A that the covariant components of the strain tensor could be expressed in terms of the components of the displacement vector $\mathbf{v}$ with respect to base vectors in B_0 or B , but these are not reproduced here since they are not used in the present work.

A stress vector $\mathbf{t}$, measured per unit area of the strained body B , and associated with the element of surface normal to the unit vector $\mathbf{n}$, can be defined at P in the strained body B , and stress quantities $-\mathbf{t}_i$ can be associated with the elementary areas at P in the surfaces $\theta_i = \text{constant}$. It was shown in A that

$$\mathbf{t} = \frac{n_i}{\sqrt{G}} \mathbf{T}_i = \tau^{ik} n_i \mathbf{E}_k = \tau_k^i n_i \mathbf{E}^k, \quad (2.8)$$

$$\mathbf{T}_i = \sqrt{(GG^{ii})} \mathbf{t}_i = \sqrt{(G)} \tau^{ik} \mathbf{E}_k = \sqrt{(G)} \tau_k^i \mathbf{E}^k, \quad (2.9)$$

$$\text{where} \quad \mathbf{n} = n_i \mathbf{E}^i, \quad G = |G_{ik}|, \quad (2.10)$$

and τ^{ik}, τ_k^i are the contravariant and mixed components of a symmetrical tensor called the stress tensor. The physical stress quantities referred to the θ_i -curves at P in B are

$$\sqrt{\left(\frac{G_{kk}}{G^{ii}}\right)} \tau^{ik}. \quad (2.11)$$

The stress equations of motion (see A) can be expressed in any of the three equivalent forms

$$\mathbf{T}_{i,i} + \rho \sqrt{(G)} \mathbf{F} = \rho \sqrt{(G)} \mathbf{f}, \quad (2.12)$$

$$\tau^{ik}{}_{||i} + \rho F^k = \rho f^k, \quad (2.13)$$

$$\tau_k^i{}_{||i} + \rho F_k = \rho f_k, \quad (2.14)$$

where ρ is the density of the strained body B , a comma denotes partial differentiation with respect to θ_i and the double line denotes covariant differentiation with respect to the strained medium B , that is, with respect to θ_i and the metric tensor components G_{ik}, G^{ik} . $\mathbf{F}$ and $\mathbf{f}$ are the body force and acceleration vectors respectively and

$$\left. \begin{aligned} \mathbf{F} &= F^k \mathbf{E}_k = F_k \mathbf{E}^k, \\ \mathbf{f} &= f^k \mathbf{E}_k = f_k \mathbf{E}^k. \end{aligned} \right\} \quad (2.15)$$

When surface forces are prescribed at a boundary,

$$\mathbf{t} = \mathbf{P}, \quad (2.16)$$

where $\mathbf{P}$ is the surface force vector measured per unit area of the deformed surface. This may be written in terms of components referred to base vectors $\mathbf{E}_k, \mathbf{E}^k$ in the forms

$$\tau^{ik} n_i = P^k, \quad \tau_k^i n_i = P_k. \quad (2.17)$$

In A it was shown that when a strain-energy function W , per unit mass, exists

$$\rho dW = \tau^{ik} d\gamma_{ik} = \rho \frac{\partial W}{\partial \gamma_{ik}} d\gamma_{ik}. \quad (2.18)$$

For a compressible elastic body the differentials $d\gamma_{ik}$ are independent, and therefore

$$\tau^{ik} = \rho \frac{\partial W}{\partial \gamma_{ik}}. \quad (2.19)$$

We may eliminate τ^{ik} from (2.13) and (2.19) to obtain differential equations for W in the form

$$\left\{ \rho \frac{\partial W}{\partial \gamma_{ik}} \right\}_{||i} + \rho F^k = \rho f^k, \quad (2.20)$$

with boundary conditions
$$\rho \frac{\partial W}{\partial \gamma_{ik}} n_i = P^k \quad (2.21)$$

at the surface of the body at which the surface forces are prescribed.

It was shown in A that Hamilton's principle for an elastic body could be deduced from the general equations. Conversely, if we assume the Hamiltonian principle, it is a simple matter to obtain the results (2.20) and (2.21) directly, but we omit the details.

This completes the summary of formulae obtained in A which will be of use in the present paper.

3. RELATIONS FOR AN INCOMPRESSIBLE BODY

When the body is incompressible, equation (2.18) is still valid, but the differentials are no longer independent. The incompressibility condition is

$$G = |G_{ik}| = g = |g_{ik}| \quad (3.1)$$

at all points of the body, so that

$$\frac{\partial G}{\partial \gamma_{ik}} d\gamma_{ik} = 0. \quad (3.2)$$

Since

$$G = |g_{ik} + 2\gamma_{ik}|,$$

it follows that

$$\frac{\partial G}{\partial \gamma_{ik}} = 2 \frac{\partial G}{\partial G_{ik}} = 2GG^{ik}, \quad (3.3)$$

G^{ik} being the contravariant components of the metric tensor of the strained body B . Using (3.3) the condition (3.2) becomes

$$G^{ik} d\gamma_{ik} = 0. \quad (3.4)$$

Hence, for an incompressible material, equation (2.18) is true for all values of the differentials $d\gamma_{ik}$ which satisfy equation (3.4); consequently the coefficients of the differentials $d\gamma_{ik}$ in the two equations must be proportional so that

$$\tau^{ik} = \rho \frac{\partial W}{\partial \gamma_{ik}} + pG^{ik}, \quad (3.5)$$

where p is an arbitrary scalar invariant function of the co-ordinates θ_i for each value of the time t . In (3.5) we consider γ_{ik} to be distinct from γ_{ki} , and the density ρ refers either to the undeformed or the deformed body.

From (2.13) and (3.5), we have, for an incompressible material,

$$\left\{ \rho \frac{\partial W}{\partial \gamma_{ik}} \right\}_{||i} + p_{,i} G^{ik} + \rho F^k = \rho f^k. \quad (3.6)$$

The corresponding boundary condition is found from (2.17) and (3.5) to be

$$\rho \frac{\partial W}{\partial \gamma_{ik}} n_i + p n^k = P^k. \quad (3.7)$$

Equations (3.6) and (3.7) may also be readily obtained from Hamilton's principle.

A number of equations given by Rivlin and earlier writers may be deduced as special cases of our general formulae. Although the deduction of these special equations has some theoretical interest, the details are not given here since the special equations given by Rivlin are much less convenient for developing the theory and application than the simple formulae (2.12) to (2.14), (2.17), together with (2.19) or (3.5), according as the body is compressible or incompressible.

4. STRAIN-ENERGY FOR ISOTROPIC MATERIALS

From now on we confine our attention to incompressible materials for which equation (3.1) holds and (3.5) may be written

$$\tau^{ik} = \frac{\partial W'}{\partial \gamma_{ik}} + p G^{ik}, \quad (4.1)$$

where W' is the strain-energy per unit volume of either the unstrained or strained body. For a material which is isotropic in the undeformed state, the strain-energy function W' is a symmetrical invariant function of the strain components γ_k^i and must therefore be a function of the three strain invariants

$$\gamma_r^r, \quad \gamma_s^r \gamma_r^s \quad \text{and} \quad |\gamma_k^i|. \quad (4.2)$$

These invariants differ from the strain invariants I_1, I_2, I_3 used by Rivlin, and it can easily be shown that

$$\left. \begin{aligned} I_1 &= 3 + 2\gamma_r^r, \\ I_2 &= 3 + 4\gamma_r^r + 2[(\gamma_r^r)^2 - \gamma_s^r \gamma_r^s], \\ I_3 &= 1 + 2\gamma_r^r + 8|\gamma_k^i| + 2[(\gamma_r^r)^2 - \gamma_s^r \gamma_r^s], \end{aligned} \right\} \quad (4.3)$$

but to conform with the notation of other writers we shall consider W' as a function of the strain-invariants I_1, I_2, I_3 instead of the invariants (4.2). When the material is incompressible, I_3 is unity and the strain-energy W' is a function of I_1 and I_2 only.

Using (4.3) and (2.6) together with the equations

$$\gamma_r^r = g^{ik} \gamma_{ik}, \quad \gamma_s^r \gamma_r^s = g^{rm} g^{sn} \gamma_{sm} \gamma_{rn}, \quad (4.4)$$

we find that

$$\frac{\partial I_1}{\partial \gamma_{ik}} = 2g^{ik}, \quad \frac{\partial I_2}{\partial \gamma_{ik}} = 2[I_1 g^{ik} - g^{ir} g^{sk} G_{rs}]. \quad (4.5)$$

Hence
$$\frac{\partial W'}{\partial \gamma_{ik}} = 2g^{ik}(\Phi + I_1 \Psi) - 2A^{ik} \Psi, \quad (4.6)$$

where we have put $A^{ik} = g^{ir} g^{sk} G_{rs}, \quad \Phi = \frac{\partial W'}{\partial I_1}, \quad \Psi = \frac{\partial W'}{\partial I_2}, \quad (4.7)$

Φ and Ψ being scalar invariant functions of I_1, I_2 .

To simplify the calculation of I_1, I_2 and I_3 in specific problems it should be noted that equations (4.3) can be written

$$I_1 = g^{rs} G_{rs}, \quad I_2 = \frac{1}{2}[I_1^2 - g^{rm} g^{sn} G_{sm} G_{rn}] = g_{rs} G^{rs} I_3, \quad I_3 = \frac{G}{g}, \quad (4.8)$$

or alternatively I_1, I_2 and I_3 can be expressed in terms of the principal extensions at a given point (see Rivlin 1948*d*).

5. PURE TORSION OF A CIRCULAR CYLINDER

This problem has already been solved by Rivlin (1948*c, d, 1949a*), but we shall outline the solution here since the final form of Rivlin's results can be obtained simply and directly from the theory developed above.

The curvilinear co-ordinate system $(\theta_1, \theta_2, \theta_3)$ in the strained state of equilibrium is taken to be a system of cylindrical polar co-ordinates (r, θ, z) , so that

$$y_1 = r \cos \theta, \quad y_2 = r \sin \theta \quad \text{and} \quad y_3 = z,$$

and the y_i -axes are supposed coincident with the x_i -axes. In the strained and unstrained states the cylinder of isotropic incompressible material occupies the region $r \leq a, 0 \leq z \leq l$, the strained state being obtained by rotating each section of the cylinder which is normal to the z -axis through an angle ψz . Thus the point (r, θ, z) in the deformed stage was originally at the point $(r, \theta - \psi z, z)$. With these assumptions we readily find that

$$\|g^{ik}\| = \begin{pmatrix} 1, & 0, & 0 \\ 0, & \frac{1}{r^2} + \psi^2, & \psi \\ 0, & \psi, & 1 \end{pmatrix}, \quad \|G^{ik}\| = \begin{pmatrix} 1, & 0, & 0 \\ 0, & \frac{1}{r^2}, & 0 \\ 0, & 0, & 1 \end{pmatrix}, \quad (5.1)$$

$$I_1 = I_2 = 3 + r^2 \psi^2, \quad (5.2)$$

and
$$\|A^{ik}\| = \begin{pmatrix} 1, & 0, & 0 \\ 0, & \frac{1}{r^2}(1 + r^2 \psi^2)^2 + \psi^2, & 2\psi + r^2 \psi^3 \\ 0, & 2\psi + r^2 \psi^3, & 1 + r^2 \psi^2 \end{pmatrix}. \quad (5.3)$$

The incompressibility condition (3.1) is satisfied throughout the body. Forming the Christoffel symbols of the second kind for the cylindrical polar co-ordinate system in the strained body and using equations (4.6) and (5.1) to (5.3) to obtain $\partial W'/\partial \gamma_{ik}$, we have, omitting the details of the algebra,

$$\left\{ \frac{\partial W'}{\partial \gamma_{ik}} \right\}_{||i} = 2 \frac{d}{dr} \{ \Phi + (2 + r^2 \psi^2) \Psi \} - 2r \psi^2 \Phi \quad (k = 1),$$

$$= 0 \quad (k = 2, 3),$$

since W' is a function of r only.

The equations of equilibrium with no body forces, obtained from (3.6), give

$$\left. \begin{aligned} 2 \frac{d}{dr} \{ \Phi + (2 + r^2 \psi^2) \Psi \} - 2r \psi^2 \Phi + \frac{\partial p}{\partial r} &= 0, \\ \frac{1}{r^2} \frac{\partial p}{\partial \theta} &= \frac{\partial p}{\partial z} = 0. \end{aligned} \right\} \quad (5.4)$$

Hence
$$p = -2 \{ \Phi + (2 + r^2 \psi^2) \Psi \} + 2 \psi^2 \int_a^r \Phi r dr, \quad (5.5)$$

and the function p is determined, apart from a constant of integration.

If we suppose that the curved surface of the cylinder is free from traction the boundary conditions (3.7) over the curved surface of the cylinder are satisfied when $k = 2$ or 3, while the remaining equation gives

$$p_{r=a} = -2 \{ \Phi + (2 + r^2 \psi^2) \Psi \}_{r=a}.$$

It now follows from (5.5) that

$$p = -2 \{ \Phi + (2 + r^2 \psi^2) \Psi \} + 2 \psi^2 \int_a^r \Phi r dr. \quad (5.6)$$

Using this value of p the boundary conditions (3.7) on the end $z = l$ of the cylinder give

and
$$\left. \begin{aligned} P^1 &= 0, & P^2 &= 2 \psi (\Phi + \Psi) \\ P^3 &= 2 \psi^2 \left\{ \int_a^r \Phi r dr - r^2 \Psi \right\} \end{aligned} \right\} \quad (5.7)$$

as the components of the surface traction.

The components of the stress tensor are given by (4.1), while the physical stress quantities are the quantities (2.11) and therefore

$$\left. \begin{aligned} \widehat{rr} &= \tau^{11} = 2 \psi^2 \int_a^r \Phi r dr, \\ \widehat{\theta\theta} &= r^2 \tau^{22} = 2 \psi^2 \left\{ r^2 \Phi + \int_a^r \Phi r dr \right\}, \\ \widehat{zz} &= \tau^{33} = 2 \psi^2 \left\{ \int_a^r \Phi r dr - r^2 \Psi \right\}, \\ \widehat{\theta z} &= r \tau^{23} = 2 \psi r \{ \Phi + \Psi \}, \\ \widehat{r\theta} &= \widehat{rz} = \tau^{12} = \tau^{13} = 0. \end{aligned} \right\} \quad (5.8)$$

From (5.8) it can be seen that when the strain-energy W' is a function of I_2 only then $\widehat{zz}$ and $\widehat{\theta z}$ are the only non-zero stress components.

The resultant couple and resultant normal force which must be applied to the end $z = l$ of the cylinder to maintain the state of strain may be determined from (5.7) or from (5.8), but we shall not derive their values here as they have been given by Rivlin (1949a).

6. ROTATION OF A RIGHT-CIRCULAR CYLINDER ABOUT ITS AXIS

As in the previous section we use cylindrical polar co-ordinates (r, θ, z) to define the points of the cylinder in the strained state, and the cylinder of incompressible isotropic material has its axis along the z -axis. The state of stress of the cylinder when it is rotated about its axis with constant angular velocity ω is identical with the stress system which would be produced in the cylinder if it were stationary under the action of a body force of magnitude $\omega^2 r$ directed along the r -axis provided that the boundary conditions remain unaltered, and we shall consider this equivalent problem to avoid the use of rotating axes.

Assuming that in the strained state there is a uniform extension λ along the z -axis so that the point (r, θ, z) was initially at the point $(r', \theta, z/\lambda)$, the incompressibility of the material implies that $r' = r\sqrt{\lambda}$, since we must have $r^2 z = r'^2 z/\lambda$. The initial co-ordinates of the point (r, θ, z) are therefore

$$x_1 = r\sqrt{\lambda} \cos \theta, \quad x_2 = r\sqrt{\lambda} \sin \theta, \quad x_3 = \frac{z}{\lambda}, \quad (6.1)$$

since the x_i - and y_i -axes are taken to coincide.

From (6.1) we find that in this case

$$\|g^{ik}\| = \begin{pmatrix} \frac{1}{\lambda}, & 0, & 0 \\ 0, & \frac{1}{\lambda r^2}, & 0 \\ 0, & 0, & \lambda^2 \end{pmatrix}, \quad (6.2)$$

while the components G^{ik} of the metric tensor of the strained body are given by (5.1). A little calculation suffices to show that

$$\|A^{ik}\| = \begin{pmatrix} \frac{1}{\lambda^2}, & 0, & 0 \\ 0, & \frac{1}{\lambda^2 r^2}, & 0 \\ 0, & 0, & \lambda^4 \end{pmatrix}, \quad (6.3)$$

$$I_1 = \lambda^2 + \frac{2}{\lambda}, \quad I_2 = 2\lambda + \frac{1}{\lambda^2}, \quad (6.4)$$

and

$$g^{ik}{}_{;i} = A^{ik}{}_{;i} = 0. \quad (6.5)$$

Since I_1 and I_2 are constants, the strain-energy W' is constant throughout the body and therefore (4.6) and (6.5) give

$$\frac{\partial W'}{\partial \gamma_{ik}}{}_{;i} = 0. \quad (6.6)$$

The equations of equilibrium are obtained by putting $F^k = (\omega^2 r, 0, 0)$ and $f^k = 0$ in (3.6), and using (6.6) they give

$$\frac{\partial p}{\partial r} = -\rho\omega^2 r, \quad \frac{1}{r^2} \frac{\partial p}{\partial \theta} = \frac{\partial p}{\partial z} = 0,$$

and it follows that

$$p = -\frac{1}{2}\rho\omega^2 r^2 + B, \quad (6.7)$$

where B is a constant to be determined.

On the curved surface of the cylinder, which we suppose to be free from traction, the unit normal is $\mathbf{n} = \mathbf{E}^1 = \mathbf{E}_1$, and (4.6), (6.2) and (6.3) show that the boundary conditions (3.7) are satisfied for $k = 2, 3$ while the other condition is

$$\frac{2}{\lambda}(\Phi + I_1\Psi) - \frac{2}{\lambda^2}\Psi - \frac{1}{2}\rho\omega^2 r^2 + B = 0 \quad \text{at} \quad r = \frac{a}{\sqrt{\lambda}}, \quad (6.8)$$

where a is the initial radius of the cylinder.

Determining the value of the constant B from (6.8) and substituting in (6.7) we obtain

$$p = \frac{1}{2}\rho\omega^2\left(\frac{a^2}{\lambda} - r^2\right) - \frac{2}{\lambda}\left\{\Phi + \left(\lambda^2 + \frac{1}{\lambda}\right)\Psi\right\}. \quad (6.9)$$

The components of the stress tensor can now be found by substituting for $\partial W'/\partial\gamma_{ik}$ and p in (4.1), and using (2.11) we find that the physical stress components are

$$\left. \begin{aligned} \widehat{rr} = \widehat{\theta\theta} &= \frac{1}{2}\rho\omega^2\left(\frac{a^2}{\lambda} - r^2\right), \\ \widehat{zz} &= \frac{1}{2}\rho\omega^2\left(\frac{a^2}{\lambda} - r^2\right) + 2\left\{\left(\lambda^2 - \frac{1}{\lambda}\right)\Phi + \left(\lambda - \frac{1}{\lambda^2}\right)\Psi\right\}, \\ \widehat{r\theta} = \widehat{\theta z} = \widehat{zr} &= 0. \end{aligned} \right\} \quad (6.10)$$

and

We see from (6.10), or alternatively from the boundary conditions (3.7), that the surface tractions on the plane ends of the cylinder consist of a normal tension of magnitude

$$T = \frac{1}{2}\rho\omega^2\left(\frac{a^2}{\lambda} - r^2\right) + 2\left\{\left(\lambda^2 - \frac{1}{\lambda}\right)\Phi + \left(\lambda - \frac{1}{\lambda^2}\right)\Psi\right\}. \quad (6.11)$$

If the traction on each end of the cylinder has no resultant then

$$\int_0^{a/\sqrt{\lambda}} Tr dr = 0,$$

or, using (6.11),
$$\left(\lambda^2 - \frac{1}{\lambda}\right)\Phi + \left(\lambda - \frac{1}{\lambda^2}\right)\Psi + \frac{1}{8}\rho\omega^2\frac{a^2}{\lambda} = 0.$$

Unless explicit expressions are assumed for Φ, Ψ , it is difficult to discuss the nature of the roots of this equation. We shall suppose that there is a unique positive root $\lambda = \lambda_1$. It is clear that, in general, this will only be possible if the angular velocity ω is not too large.

Substituting $\lambda = \lambda_1$, in (6.10) we find that the stresses are

$$\left. \begin{aligned} \widehat{rr} = \widehat{\theta\theta} &= \frac{1}{2}\rho\omega^2\left(\frac{a^2}{\lambda_1} - r^2\right), \\ \widehat{zz} &= \frac{1}{4}\rho\omega^2\left(\frac{a^2}{\lambda_1} - 2r^2\right). \end{aligned} \right\} \quad (6.12)$$

When the strain is infinitesimal λ_1 is slightly less than unity and, to the order of approximation of the infinitesimal theory, we may put $\lambda_1 = 1$ in (6.12), obtaining the classical results for the special case when the material is incompressible.

7. SPHERICAL SHELL OF INCOMPRESSIBLE MATERIAL

It is convenient to identify the curvilinear co-ordinate system in the strained body with a system of spherical polar co-ordinates which has its origin at the centre of the shell. Then writing R, θ, ϕ for $\theta_1, \theta_2, \theta_3$ we have

$$y_1 = R \sin \theta \cos \phi, \quad y_2 = R \sin \theta \sin \phi, \quad y_3 = R \cos \theta, \quad (7.1)$$

and

$$\|G^{ik}\| = \begin{pmatrix} 1, & 0, & 0 \\ 0, & \frac{1}{R^2}, & 0 \\ 0, & 0, & \frac{1}{R^2 \sin^2 \theta} \end{pmatrix}, \quad G = R^4 \sin^2 \theta. \quad (7.2)$$

We assume that the displacement of the unstrained body possesses spherical symmetry so that the point (R, θ, ϕ) was originally at the point (r, θ, ϕ) and we shall write

$$Q(R) = \frac{r}{R} \quad (7.3)$$

for convenience in the algebra. The internal and external radii of the shell in the unstrained and strained states will be denoted by r_1, r_2 and R_1, R_2 respectively. For incompressibility we must have

$$r^3 - R^3 = r_1^3 - R_1^3 = r_2^3 - R_2^3 \quad (7.4)$$

and therefore

$$Q(R) = \left(1 + \frac{r_1^3 - R_1^3}{R^3}\right)^{\frac{1}{3}}. \quad (7.5)$$

Taking the x_i -axes to coincide with the y_i -axes we have, using (7.3),

$$x_1 = RQ \sin \theta \cos \phi, \quad x_2 = RQ \sin \theta \sin \phi, \quad x_3 = RQ \cos \theta. \quad (7.6)$$

From (7.5) we deduce that

$$\frac{dQ}{dR} = \frac{1}{R} \left(\frac{1}{Q^2} - Q \right), \quad (7.7)$$

and using this result and (7.6) we find that

$$\|g^{ik}\| = \begin{pmatrix} Q^4, & 0, & 0 \\ 0, & \frac{1}{R^2 Q^2}, & 0 \\ 0, & 0, & \frac{1}{R^2 Q^2 \sin^2 \theta} \end{pmatrix}, \quad g = R^4 \sin^2 \theta. \quad (7.8)$$

We also have

$$\|A^{ik}\| = \begin{pmatrix} Q^8, & 0, & 0 \\ 0, & \frac{1}{R^2 Q^4}, & 0 \\ 0, & 0, & \frac{1}{R^2 Q^4 \sin^2 \theta} \end{pmatrix}, \quad (7.9)$$

$$I_1 = Q^4 + \frac{2}{Q^2} \quad \text{and} \quad I_2 = 2Q^2 + \frac{1}{Q^4}. \quad (7.10)$$

The non-zero Christoffel symbols of the second kind for the strained body are

$$\begin{aligned} \left\{ \begin{matrix} 1 \\ 22 \end{matrix} \right\} &= -R, & \left\{ \begin{matrix} 1 \\ 33 \end{matrix} \right\} &= -R \sin^2 \theta, & \left\{ \begin{matrix} 2 \\ 33 \end{matrix} \right\} &= -\sin \theta \cos \theta, \\ \left\{ \begin{matrix} 2 \\ 12 \end{matrix} \right\} &= \left\{ \begin{matrix} 3 \\ 13 \end{matrix} \right\} = \frac{1}{R} & \text{and} & \left\{ \begin{matrix} 3 \\ 23 \end{matrix} \right\} = \frac{\cos \theta}{\sin \theta}, \end{aligned}$$

and omitting the details of the calculation, we find from (4.6), (7.7) to (7.10) that

$$\begin{aligned} \frac{\partial W'}{\partial \gamma_{ik}} \parallel_i &= \frac{2}{R} \left(\frac{1}{Q^2} - Q \right) \left\{ Q^4 \frac{d\Phi}{dQ} + 2Q^2 \frac{d\Psi}{dQ} + 2(Q^3 - 1) \Phi + 2 \left(Q - \frac{1}{Q^2} \right) \Psi \right\} \quad (k = 1), \\ &= 0 \quad (k = 2, 3). \end{aligned} \quad (7.11)$$

The equations of motion (3.6) with no body forces and no acceleration are therefore

$$\begin{aligned} \frac{2}{R} \left(\frac{1}{Q^2} - Q \right) \left\{ Q^4 \frac{d\Phi}{dQ} + 2Q^2 \frac{d\Psi}{dQ} + 2(Q^3 - 1) \Phi + 2 \left(Q - \frac{1}{Q^2} \right) \Psi \right\} + \frac{\partial p}{\partial R} &= 0, \\ \frac{1}{R^2} \frac{\partial p}{\partial \theta} &= \frac{1}{R^2 \sin^2 \theta} \frac{\partial p}{\partial \phi} = 0. \end{aligned} \quad (7.12)$$

The last two equations show that p is a function of R only, and from (7.7) and (7.12) we have

$$\frac{dp}{dQ} = -2 \left\{ Q^4 \frac{d\Phi}{dQ} + 2Q^2 \frac{d\Psi}{dQ} + 2(Q^3 - 1) \Phi + 2 \left(Q - \frac{1}{Q^2} \right) \Psi \right\}.$$

Integrating and using the rule of integration by parts we obtain

$$p = -2Q^2(Q^2\Phi + 2\Psi) + 4 \int^Q \left\{ (Q^3 + 1) \Phi + \left(Q + \frac{1}{Q^2} \right) \Psi \right\} dQ, \quad (7.13)$$

and the function p is determined apart from a constant of integration.

If Π_1 is the normal pressure on the inner surface of the shell, the unit normal to which is $\mathbf{n} = -\mathbf{E}_1 = -\mathbf{E}^1$, the boundary conditions (3.7) on this surface show that

$$P^1 = \Pi_1 = -p - \frac{\partial W'}{\partial \gamma_{11}},$$

and

$$P^2 = P^3 = 0 \quad \text{at} \quad R = R_1,$$

and using the first of these equations together with (4.6), (7.8) to (7.10) and (7.13) we find that

$$p = -2Q^2\{Q^2\Phi + 2\Psi\} + K(R) - \Pi_1, \quad (7.14)$$

$$\text{where} \quad K(R) = 4 \int_{Q_1}^Q \left\{ (Q^3 + 1) \Phi + \left(Q + \frac{1}{Q^2} \right) \Psi \right\} dQ, \quad Q_1 = \frac{r_1}{R_1}. \quad (7.15)$$

The stresses can now be determined from (4.1) and (2.11), and we obtain

$$\begin{aligned} \widehat{RR} &= \tau^{11} = K(R) - \Pi_1, \\ \widehat{\theta\theta} &= \widehat{\phi\phi} = R^2 \tau^{22} = R^2 \sin^2 \theta \tau^{33} = 2 \left\{ \left(\frac{1}{Q^2} - Q^4 \right) \Phi + \left(\frac{1}{Q^4} - Q^2 \right) \Psi \right\} + K(R) - \Pi_1, \\ \widehat{\theta\phi} &= \widehat{\phi R} = \widehat{R\theta} = 0. \end{aligned} \quad (7.16)$$

The boundary conditions (3.7), or more simply the equations (7.16) and the boundary conditions (2.17), show that if Π_2 is the normal pressure on the external surface of the shell then

$$\Pi_2 = \Pi_1 - K(R_2). \quad (7.17)$$

Since we can express R_2 in terms of R_1 by the incompressibility condition (7.4), this equation serves to determine R_1 if the difference in pressure $\Pi_1 - \Pi_2$ and the form of the strain-energy function W' are known; alternatively, it gives the value of $\Pi_1 - \Pi_2$ when R_1 is known.

It may happen that the outer radius of the shell is infinite and the body is then an infinite solid containing a spherical cavity. In this case a possible form of the boundary conditions at infinity is that the stresses should vanish there, and noting that $Q(R) \rightarrow 1$ as $R \rightarrow \infty$ we see from (7.16) that this condition will be satisfied if

$$\Pi_1 = K(\infty),$$

provided Φ and Ψ are finite at infinity.

The expression (7.15) for $K(R)$ can be put into a different form. From (7.10) we have

$$\frac{dI_1}{dQ} = 4 \frac{(Q^3 - 1)}{Q^3} (Q^3 + 1), \quad \frac{dI_2}{dQ} = 4 \frac{(Q^3 - 1)}{Q^3} \left(Q + \frac{1}{Q^2} \right),$$

and therefore

$$\frac{dW'}{dQ} = 4 \frac{(Q^3 - 1)}{Q^3} \left\{ (Q^3 + 1) \Phi + \left(Q + \frac{1}{Q^2} \right) \Psi \right\}.$$

Hence, using (7.5), (7.15) can be written

$$K(R) = \int_{Q_1}^Q \frac{Q^3}{Q^3 - 1} \frac{dW'}{dQ} dQ = \int_{R_1}^R \frac{(R^3 + r_1^3 - R_1^3)}{r_1^3 - R_1^3} \frac{dW'}{dR} dR. \quad (7.18)$$

Putting $R = R_2$ in (7.18), integrating by parts and remembering (7.4) we find that

$$K(R_2) = \frac{1}{r_1^3 - R_1^3} \left\{ r_2^3 W'(R_2) - r_1^3 W'(R_1) - \frac{3}{4\pi} \bar{W} \right\}, \quad (7.19)$$

where $\bar{W} = \int_{R_1}^{R_2} 4\pi R^2 W' dR$ is the total energy stored elastically by the body. Using (7.4) again, (7.17) and (7.19) give

$$\frac{4}{3}\pi(R_1^3 - r_1^3)\Pi_1 - \frac{4}{3}\pi(R_2^3 - r_2^3)\Pi_2 = \bar{W} + \frac{4}{3}\pi r_1^3 W'(R_1) - \frac{4}{3}\pi r_2^3 W'(R_2). \quad (7.20)$$

Equations (7.16) and (7.17) can be used to derive the approximate formulae for the stresses when the strain is infinitesimal, and it is found that their values agree with those obtained from the classical theory.

We shall now apply the results deduced above to obtain the corresponding results for a thin spherical shell. In this case we put $r_2 = r_1(1 + \epsilon)$, where ϵ is small enough to allow us to neglect second degree and higher terms in ϵ . Denoting the extension of the shell by λ , so that $R_1 = \lambda r_1$, we have from (7.4)

$$R_2 = \lambda r_1 \left(1 + \frac{\epsilon}{\lambda^3} \right)$$

approximately, λ being considered large compared with ϵ . It follows that

$$Q_1 = \frac{r_1}{R_1} = \frac{1}{\lambda}, \quad Q_2 = \frac{r_2}{R_2} = \frac{1}{\lambda} \left[1 + \epsilon \left(1 - \frac{1}{\lambda^3} \right) \right], \quad (7.21)$$

and we see that $Q_2 - Q_1$ is of the same order as ϵ .

Hence, to the first order in ϵ , we have from (7.15)

$$K(R_2) = 4(Q_2 - Q_1) \left[(Q^3 + 1) \Phi + \left(Q + \frac{1}{Q^2} \right) \Psi \right]_{R=R_1}.$$

Introducing the values (7.21) of Q_1 and Q_2 into this equation and using (7.17) we find that

$$\Pi_1 - \Pi_2 = 4\epsilon \left\{ \left(\frac{1}{\lambda} - \frac{1}{\lambda^7} \right) (\Phi)_{R=R_1} + \left(\lambda - \frac{1}{\lambda^5} \right) (\Psi)_{R=R_1} \right\}. \quad (7.22)$$

For the neo-Hookean incompressible material defined by Rivlin (1948*a*) we have $W' = C_1(I_1 - 3)$, where C_1 is a constant, and (7.22) becomes

$$\Pi_1 - \Pi_2 = 4\epsilon C_1 \left(\frac{1}{\lambda} - \frac{1}{\lambda^7} \right). \quad (7.23)$$

We see that, for this material, $\Pi_1 - \Pi_2$ increases to a maximum as λ increases from unity to $\sqrt[6]{7} = 1.383$, and then decreases to zero as λ increases to infinity. As λ decreases from unity to zero $\Pi_1 - \Pi_2$ decreases monotonically from zero to negative infinity.

When the material is such that the strain-energy $W' = C_2(I_2 - 3)$, where C_2 is a constant, (7.22) becomes

$$\Pi_1 - \Pi_2 = 4\epsilon C_2 \left(\lambda - \frac{1}{\lambda^5} \right),$$

and for this material $\Pi_1 - \Pi_2$ is a monotonic increasing function of λ .

One of the writers (R.T.S.) wishes to thank the Department of Scientific and Industrial Research for financial assistance.

REFERENCES

- Green, A. E. & Zerna, W. 1950 *Phil. Mag.* **41**, 313.
Murnaghan, F. D. 1937 *Amer. J. Math.* **59**, 235.
Rivlin, R. S. 1948*a* *Phil. Trans. A*, **240**, 459.
Rivlin, R. S. 1948*b* *Phil. Trans. A*, **240**, 491.
Rivlin, R. S. 1948*c* *Phil. Trans. A*, **240**, 509.
Rivlin, R. S. 1948*d* *Phil. Trans. A*, **241**, 379.
Rivlin, R. S. 1949*a* *Proc. Camb. Phil. Soc.* **45**, 485.
Rivlin, R. S. 1949*b* *Proc. Roy. Soc. A*, **195**, 463.
Rivlin, R. S. 1949*c* *Phil. Trans. A*, **242**, 173.

The quaternary system aluminium-iron-cobalt-nickel, with reference to the role of transitional metals in alloys

By G. V. RAYNOR AND M. B. WALDRON

(Communicated by A. J. Bradley, F.R.S.—Received 1 March 1950)

In a previous publication, the results of an investigation of the aluminium-rich portion of the aluminium-iron-nickel equilibrium diagram were reported and interpreted theoretically (Raynor & Pfeil 1946-7*a*). On the basis of this theoretical work, the forms of the previously unknown equilibrium diagrams for the systems aluminium-iron-cobalt and aluminium-cobalt-nickel were deduced, and subsequently verified quantitatively by experiment (Raynor & Pfeil 1946-7*b*; Raynor & Waldron 1948). In the present paper, similar reasoning has been applied to the quaternary aluminium-iron-cobalt-nickel system, and the form of diagram expected compared with experiment. Previous work has suggested that Co_2Al_9 and the isomorphous FeNiAl_9 are analogous to electron compounds. If the theory be accepted that transitional metal atoms may absorb electrons in aluminium-rich alloys to an extent governed by the vacancies in their atomic orbitals (or $3d$ shells), the two compounds have closely similar electron:atom ratios, and both dissolve nickel to the same limiting electron:atom ratio. Co_2Al_9 will dissolve more iron than FeNiAl_9 ; the reason for this has been discussed. In the quaternary system, it would be expected that a continuous series of solid solutions, of composition $(\text{Fe} \cdot \text{Co} \cdot \text{Ni})_2\text{Al}_9$, would exist between the two compounds, and that the aluminium-rich boundary of the quaternary body in the tetrahedral equilibrium model would be a plane corresponding to a constant proportion of 2 solute atoms to 9 aluminium atoms. It would also be expected that the whole boundary corresponding to saturation of the quaternary body with nickel would occur at a constant electron:atom ratio. These considerations imply that no other phases, apart from FeAl_3 and NiAl_3 (which dissolve relatively small amounts of the other transitional solutes), enter into equilibrium with the aluminium-rich solid solution, and that the liquidus surfaces for FeNiAl_9 and Co_2Al_9 merge continuously into each other without a break. Further predictions with regard to the form of the equilibrium diagram may also be made, and these are discussed in the paper.

The results of the investigation, which are discussed, confirm previous suggestions that the inter-metallic compounds formed by aluminium and transitional metals may, if sufficient aluminium is present, be regarded as analogous to electron compounds, with absorption of electrons by the transitional metal atoms occurring. The theory may be used, in favourable cases, to predict quantitatively the forms of uninvestigated complex equilibrium diagrams. In the present case, equilibrium relationships in a quaternary system have been predicted from those in the three subsidiary ternary systems, two of which were themselves largely predicted as a result of the original theoretical interpretation of the aluminium-iron-nickel alloys.

1. INTRODUCTION

Investigations of the constitutions of the ternary systems aluminium-iron-nickel, aluminium-cobalt-nickel and aluminium-iron-cobalt have already been reported and discussed (Raynor & Pfeil 1946-7*a, b*; Raynor & Waldron 1948). The main conclusions from this work, which was undertaken to obtain information concerning the theory of transitional metals in alloys, were, first, that in aluminium-rich alloys transitional metal atoms may accept electrons from the structure as a whole, and, secondly, that the binary and ternary compounds occurring in the three ternary systems may be regarded as analogous in several respects to electron compounds. These researches, taken in conjunction with other work, suggest that the numbers of electrons per atom accepted in aluminium-rich alloys by transitional metals of the first long period are closely similar to the vacancies per atom in the 'atomic orbitals'

postulated in the Pauling theory of these metals (Pauling 1938). For iron, cobalt and nickel these vacancies per atom are respectively 2.66, 1.71 and 0.61. Thus, from consideration of the recently determined crystal structure of the phase Co_2Al_9 (Douglas 1948), it has been shown that the inscribed sphere to the most probable first Brillouin zone corresponds to approximately 2.12 electrons per atom, which is consistent with the acceptance of 1.7 electrons per atom by cobalt (Douglas 1948; Raynor & Waldron 1949). Similar conclusions follow from a study of the structure of Co_2Al_5 (A. M. B. Douglas, private communication), while the form of the most probable first Brillouin zone for NiAl_3 is consistent with the acceptance of 0.6 electron per nickel atom (Raynor & Waldron 1948).

The mechanism of the apparent acceptance of a non-integral number of electrons per atom is obscure; it would be expected, according to the conventional band theory of the transitional metals, that iron, cobalt and nickel might respectively accept, into vacant levels in the $3d$ band, 2, 1 and 0 electrons per atom. The experimental evidence at present available appears to justify the retention of the non-integral values, though it will be appreciated that the alternative use of the integral values would not affect the principles of the theory of transitional metals in aluminium alloys, but merely the magnitudes of the electron:atom ratios calculated for the various phases. In the present case, the relative magnitudes of the electron:atom ratios are unaffected, since a range of compositions corresponding to a constant electron:atom ratio using the non-integral values also corresponds to an effectively constant ratio using the integral values. It seemed of interest to examine how adequately the constitution of the hitherto unexamined aluminium-rich aluminium-iron-cobalt-nickel alloys could be deduced from a knowledge of the three relevant ternary systems, in the same way as the main features of the aluminium-cobalt-nickel and aluminium-iron-cobalt equilibrium diagrams were deduced from a knowledge of the constitution of the aluminium-iron-nickel alloys (Raynor & Pfeil 1946-7*b*; Raynor & Waldron 1948). The present paper describes experiments carried out to test the validity of the conclusions reached. For brevity, the ternary systems involved will be referred to in terms of their component metals (e.g. Al-Fe-Ni, Al-Co-Ni, Al-Fe-Co), while the quaternary system will be referred to as Al-Fe-Co-Ni.

2. THE TERNARY SYSTEMS AND THE PROBABLE FORM OF THE QUATERNARY EQUILIBRIUM DIAGRAM

(i) *The system Al-Fe-Ni*

This system was examined by Bradley & Taylor (1940), by Schrader & Hanemann (1943), and more recently by Raynor & Pfeil (1946-7*a*), while the liquidus surfaces have been determined by Phillips (1942). Only features essential to the understanding of the quaternary system will be described here. The projection of the surfaces of primary separation, and the constitution in the solid state at 600°C , are shown respectively in figure 1*a* and *b*. The ternary compound $T(\text{FeNi})$, which is formed peritectically from FeAl_3 , includes the composition FeNiAl_9 and is isomorphous with Co_2Al_9 . Along the aluminium-rich boundary of the ternary phase, which involves equilibrium with the primary aluminium-rich solid solution α , nickel

and iron replace each other atom for atom. The phase may therefore be denoted $(\text{Fe.Ni})_2\text{Al}_3$, the relative proportions of nickel and iron being variable. According to the theory outlined above, the homogeneity range of the ternary phase corresponds

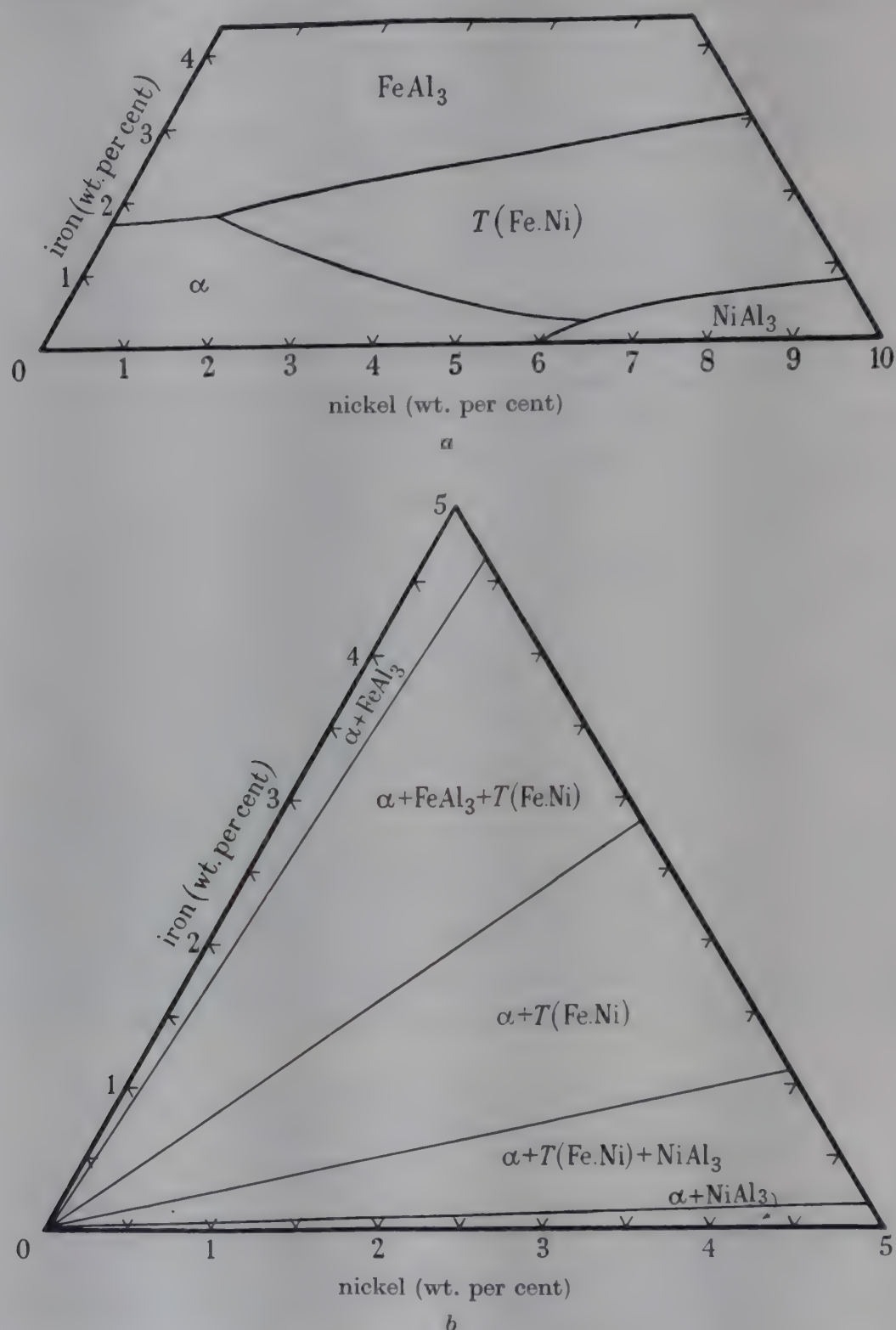


FIGURE 1. The system Al-Fe-Ni; *a*, projection of surfaces of primary separation. *b*, constitution at 600° C.

with electron : atom ratios within the extreme limits 2.14 and 2.28. The upper limit was deduced from the examination of extracted primary crystals; the corresponding limits at 600°C, as determined from the isothermal diagram, are 2.14 and 2.26

electrons per atom. FeAl_3 dissolves nickel by replacement of iron atoms until an electron : atom ratio of 1.76 is reached, while NiAl_3 dissolves a small amount (approximately 1.2 %) of iron. In the 600° C isothermal diagram, therefore, the $(\alpha + \text{NiAl}_3)$ and $(\alpha + \text{FeAl}_3)$ fields are narrow, while the $(\alpha + T(\text{FeNi}))$ field is comparatively wide.

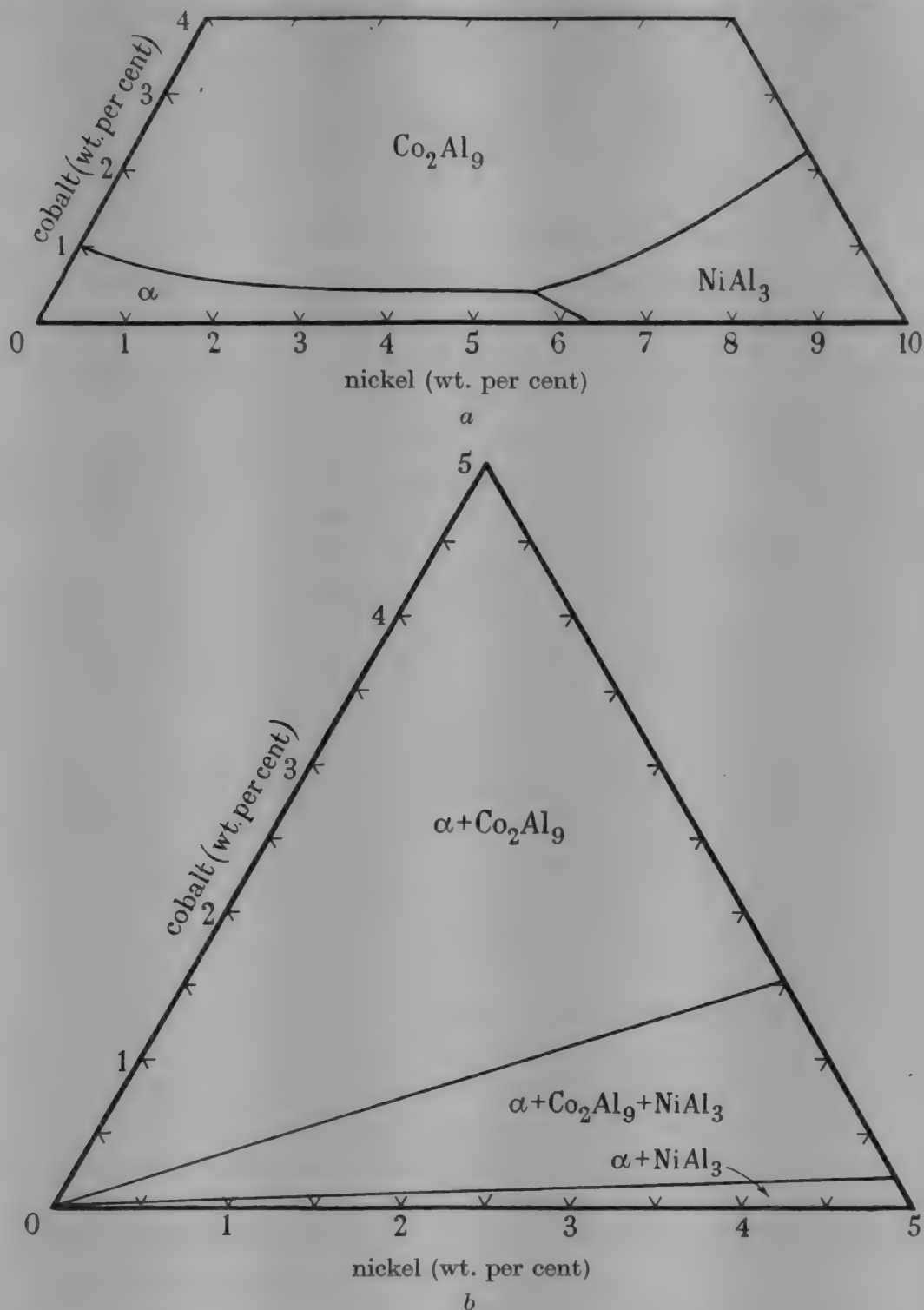


FIGURE 2. The system Al-Co-Ni; *a*, projection of surfaces of primary separation. *b*, constitution at 600° C.

(ii) The system Al-Co-Ni

From a knowledge of the Al-Fe-Ni system, theoretical predictions of the form of the Al-Co-Ni constitutional diagram may be made; these have been tested by Raynor & Pfeil (1946-7*b*). In this system no ternary compound exists in equilibrium

with the primary solid solution. Equilibrium in the aluminium-rich regions involves only the α , Co_2Al_9 and NiAl_3 phases, as shown in figure 2*a* and *b*, which represents respectively the projection of the liquidus surfaces and the constitution at 600°C . The $(\alpha + \text{Co}_2\text{Al}_9)$ field at 600°C is extensive, indicating a large solubility of nickel in Co_2Al_9 . This solubility extends to an electron:atom ratio very close to 2.28, in excellent agreement with the value at the solubility limit of nickel in $(\text{Fe.Ni})_2\text{Al}_9$. At 600°C , the appropriate electron:atom ratio is again 2.26. NiAl_3 dissolves approximately 2.8 % cobalt by replacement of nickel atoms.

(iii) *The system Al-Fe-Co*

This system was examined by Raynor & Waldron (1948) in a further investigation of the theoretical principles involved. The projection of the liquidus surfaces, and the constitution at 600°C , are shown in figure 3*a* and *b*. No ternary compound exists in equilibrium with the primary solid solution; there is a wide $(\alpha + \text{Co}_2\text{Al}_9)$ field and a comparatively narrow $(\alpha + \text{FeAl}_3)$ region. The solubility of cobalt in FeAl_3 extends to an electron:atom ratio of 1.75, in close agreement with the solubility of nickel in FeAl_3 . The limiting solubility of iron in Co_2Al_9 corresponds with an electron:atom ratio of 2.06, or, at 600°C , 2.07.

The relationships between the three ternary systems have been fully discussed (Raynor & Waldron 1948); there is a close quantitative correspondence between the Al-Co-Ni system and the portion of the Al-Fe-Ni system lying to the nickel-rich side of the line representing equal proportions of iron and nickel. There is also close correspondence between the Al-Fe-Co alloys and the portion of the Al-Fe-Ni system lying to the iron-rich side of the equal solute concentration line. The aluminium-rich Al-Fe-Ni equilibrium diagram may, in fact, be regarded as built up from the other two ternary systems placed side by side with the Al-Co axis common to both. For the present purpose, the following features are of particular interest:

(a) At 600°C , FeAl_3 dissolves both cobalt and nickel to the same limiting electron:atom ratio of 1.75 to 1.76.

(b) The solubility of both iron and cobalt in NiAl_3 is limited.

(c) The isomorphous phases $(\text{Fe.Ni})_2\text{Al}_9$ and Co_2Al_9 both dissolve nickel to a limiting electron:atom ratio of approximately 2.28. This suggests that the first Brillouin zone for the Co_2Al_9 structure may be filled up to 2.28 electrons per atom before the compound becomes unstable relative to possible alternatives; the recent work on the Brillouin zone structure (Douglas 1948, Raynor & Waldron 1949) is consistent with this. Owing to the effect of temperature, the corresponding electron:atom ratios representing the solubilities of nickel in Co_2Al_9 and $(\text{Fe.Ni})_2\text{Al}_9$ at 600°C are somewhat smaller (2.26 in each case).

(d) Co_2Al_9 withstands a deficit of electrons, and dissolves iron until the electron:atom ratio is reduced to 2.06 (2.07 at 600°C). $(\text{Fe.Ni})_2\text{Al}_9$, however, dissolves iron only until the electron:atom ratio is reduced to 2.14. The reason for this difference between the two isomorphs has been previously discussed (Raynor & Waldron 1948), and it is suggested that $(\text{Fe.Ni})_2\text{Al}_9$ is intrinsically less stable than Co_2Al_9 . Co_2Al_9 , in which all the transitional metal sites are occupied by cobalt atoms,

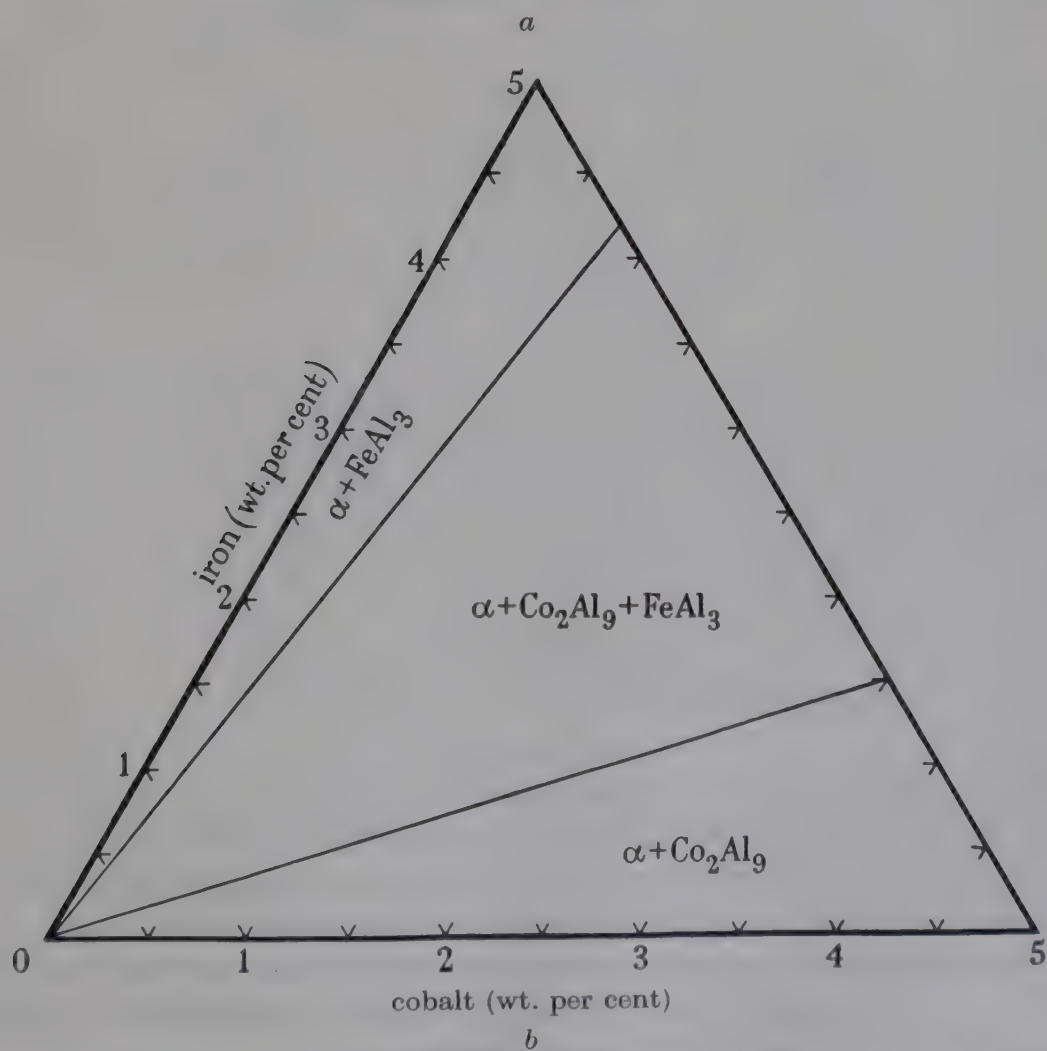
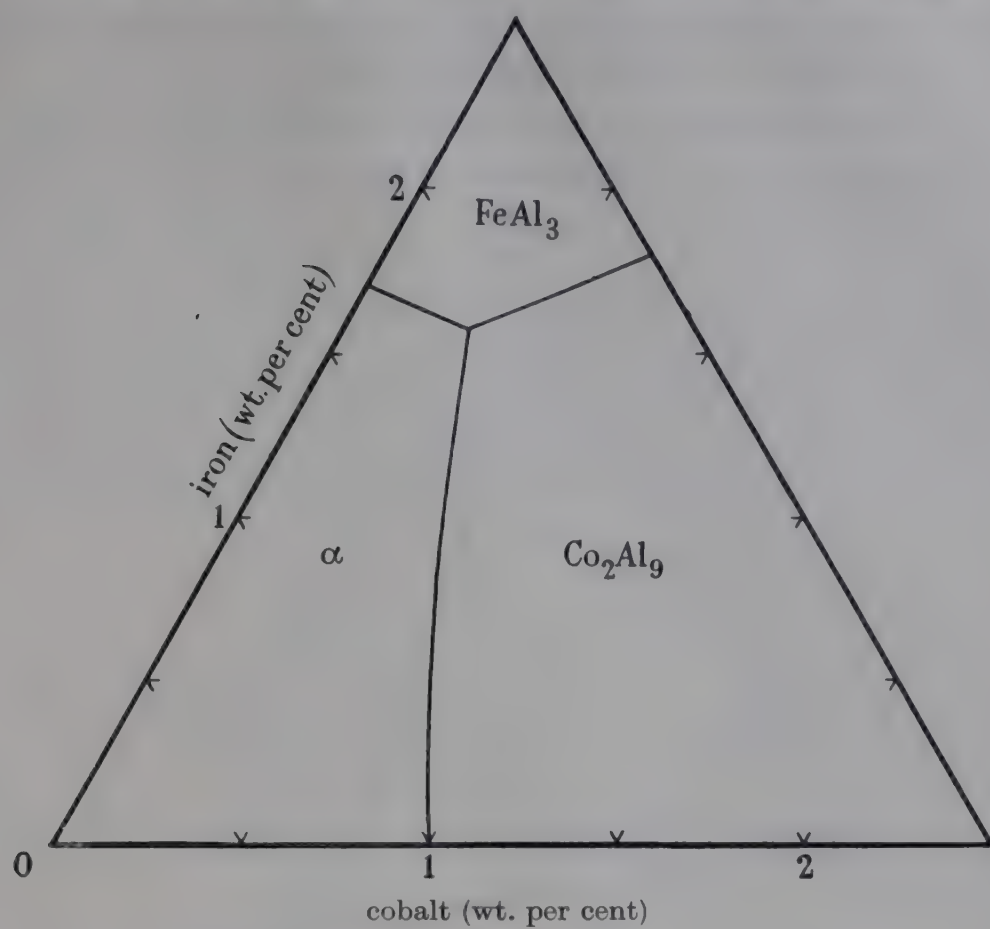


FIGURE 3. The system Al-Fe-Co; *a*, projection of surfaces of primary separation. *b*, constitution at 600° C.

enters into a eutectic relationship with FeAl_3 , and is stable up to 943°C . $(\text{Fe.Ni})_2\text{Al}_9$, however, is formed peritectically from FeAl_3 at 809°C .

The probable general form of the Al-Fe-Co-Ni equilibrium diagram at the aluminium-rich corner of the tetrahedral model may now be deduced. Since

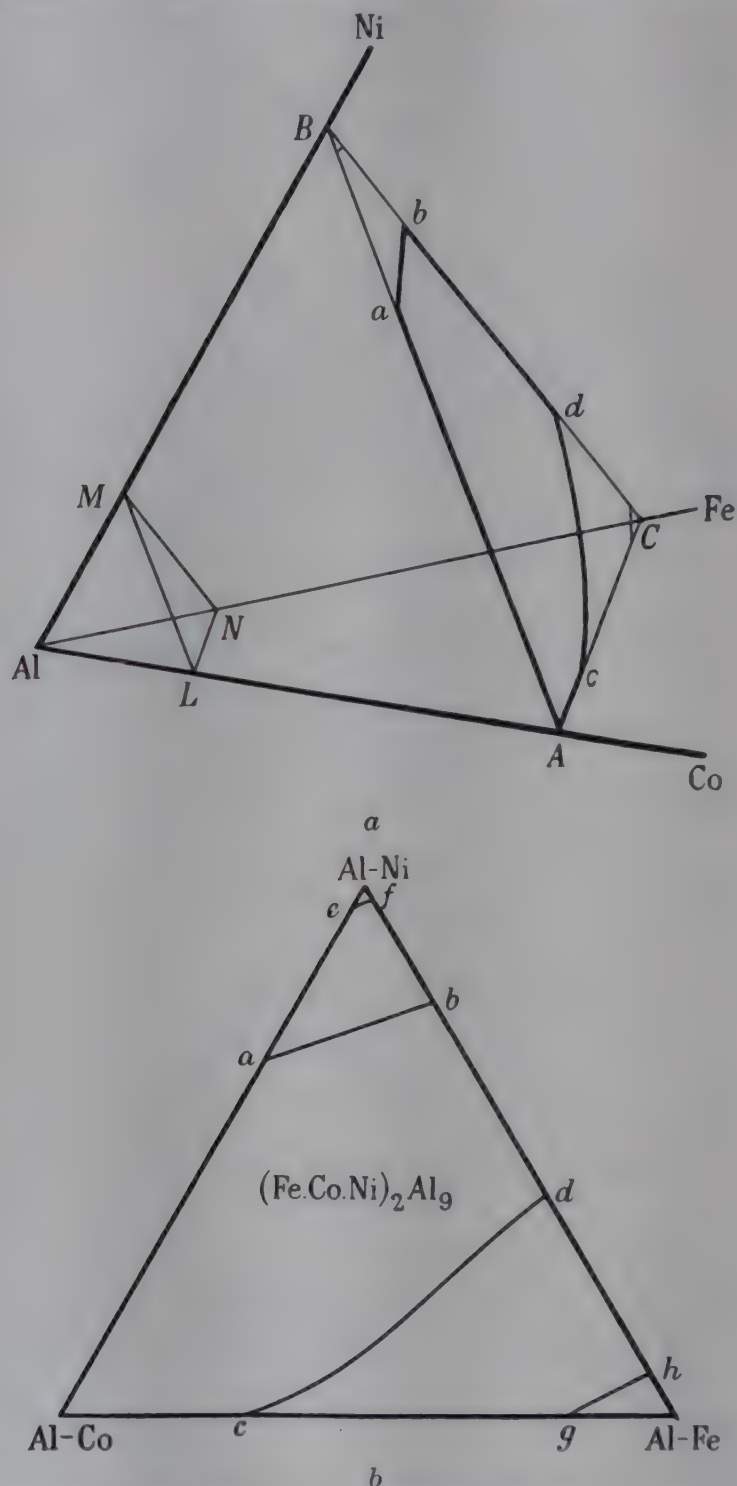


FIGURE 4. *a*, tetrahedral model of the system Al-Fe-Co-Ni. Plane ABC represents a constant ratio of 9 Al atoms to 2 solute atoms. *b*, the plane ABC in plan.

$(\text{Fe.Ni})_2\text{Al}_9$ and Co_2Al_9 have the same crystal structure, and since both dissolve iron and nickel, which differ only slightly in atomic size from cobalt, it may be expected that a continuous series of solid solutions would exist between them in the quaternary alloys. No other phase, apart from FeAl_3 and NiAl_3 , would therefore be expected to enter into equilibrium with the α solid solution.

These considerations imply that, in the tetrahedral model, the aluminium-rich boundary of the quaternary body, which may be denoted $(\text{Fe.Co.Ni})_2\text{Al}_9$, lies in a plane of constant aluminium content corresponding to 2 solute atoms per 9 aluminium atoms. The appropriate plane, ABC , is shown in figure 4*a*, which represents a perspective view of the tetrahedral model plotted in terms of atomic percentages. Since both $(\text{Fe.Ni})_2\text{Al}_9$ and Co_2Al_9 dissolve nickel to the same electron : atom ratio, the limiting solubility of nickel in the quaternary body should be represented, in the plane ABC , by a straight line of constant electron : atom ratio (2.28). Owing to the effect of temperature, the corresponding straight line at 600°C would be expected to occur at 2.26 electrons/atom, as in the component ternary systems. The solubility of iron at 600°C would be represented by a curve joining the

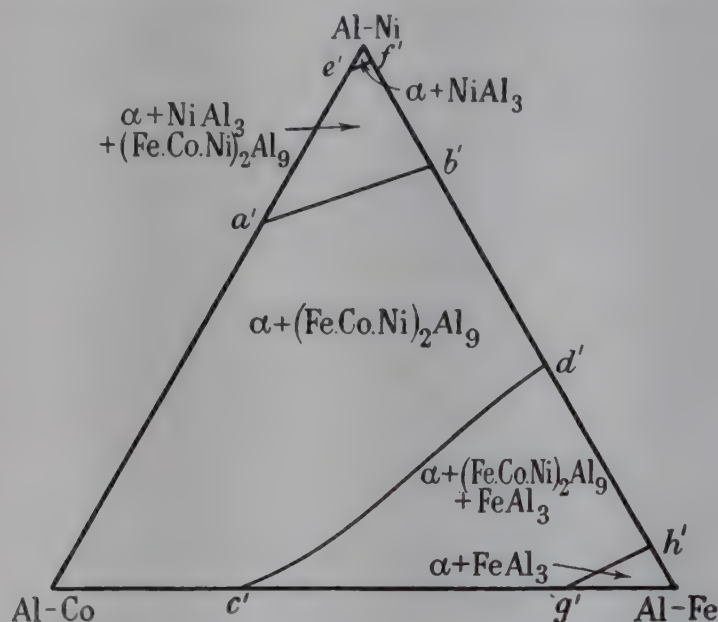


FIGURE 5. Form of equilibrium relations expected, in Al-rich alloys, for a plane of constant Al content.

composition corresponding to 2.07 electrons per atom on the Al-Fe-Co face of the tetrahedron to the composition corresponding to 2.14 electrons per atom on the Al-Fe-Ni face. The solubility of cobalt and nickel in FeAl_3 should be given by a straight line of electron : atom ratio 1.75 to 1.76. Thus the aluminium-rich boundary of the homogeneity range of the quaternary body may be represented as defined by the lines ab and cd on the plane ABC , which is shown in plan in figure 4*b*. This figure also contains lines ef and gh which respectively bound the $(\alpha + \text{NiAl}_3)$ and $(\alpha + \text{FeAl}_3)$ fields.

The solubilities of iron, cobalt and nickel in solid aluminium are negligible. The various phase boundaries separating two-phase and three-phase fields at a given temperature in the quaternary model thus diverge from the aluminium-rich apex as ruled surfaces; there should be, according to the present theory, no four-phase fields in the aluminium-rich isothermals. From figure 4*a* it will be appreciated that any plane of constant aluminium content, such as LMN , lying between the aluminium-rich apex and ABC , is a projection of the latter plane, so that the phase fields are of the same geometrical form. The field labelled $(\text{Fe.Co.Ni})_2\text{Al}_9$ in figure 4*b* becomes an $(\alpha + (\text{Fe.Co.Ni})_2\text{Al}_9)$ field, as shown in figure 5. If the ideas of this section are

correct, any aluminium-rich section of the quaternary isothermal at 600°C and constant aluminium content should be of the form of figure 5.

Similar, but less quantitative, suggestions are possible concerning the form of the liquidus surfaces for the quaternary alloys. Projections of the regions of primary separation for the three ternary systems are shown in figures 1*a*, 2*a* and 3*a*. Since Co_2Al_9 and $(\text{Fe.Ni})_2\text{Al}_9$ are expected to merge into each other, the field of primary separation of Co_2Al_9 would be expected to merge continuously into that of the phase $(\text{Fe.Ni})_2\text{Al}_9$. The approximate forms of the composition regions within which the various phases in the quaternary system separate as primary constituents may therefore be represented as shown in figure 6.

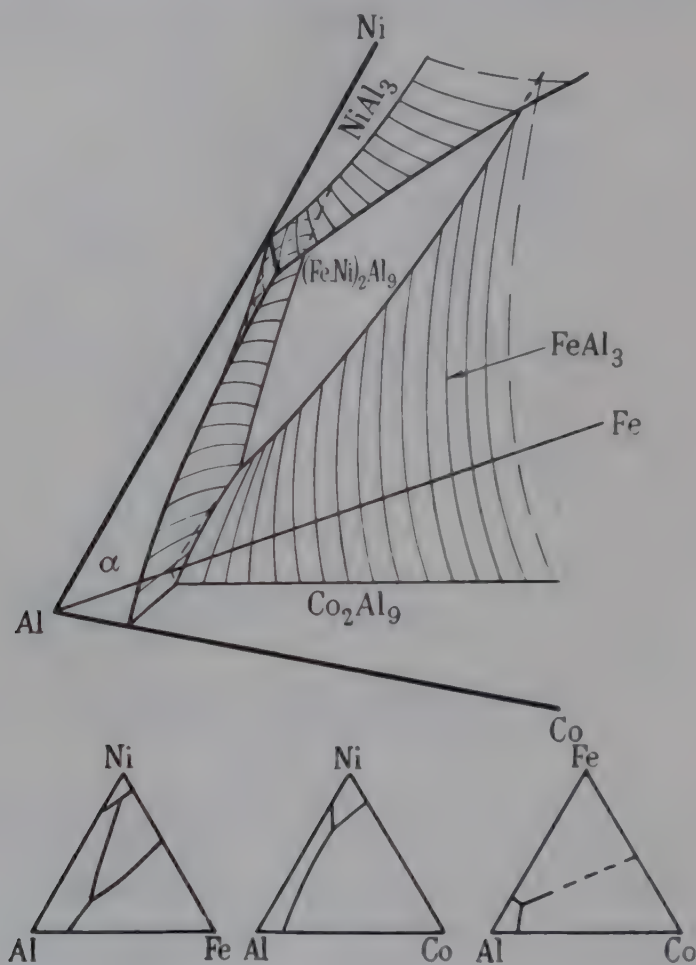


FIGURE 6. Expected form of the liquidus surfaces.

The experiments described below were undertaken to test the validity of these suggestions.

3. CONSTITUTION IN THE SOLID STATE; THE 600°C ISOTHERMAL

The most accurate isothermal diagrams for the three ternary systems were established at 600°C . It was considered, therefore, that a quaternary isothermal at 600°C would provide the most direct check on the suggestions made above. The experimental methods used were closely similar to those described previously (Raynor & Waldron 1948). For convenience, the nominal compositions of the alloys prepared lay in the plane of constant aluminium content at 97.5% aluminium by weight; the suggestions of the previous section enabled the important features of the

isothermal to be examined with a relatively small number of alloys. Specimens were annealed, quenched in water and examined micrographically.

During the course of the work, no phases other than FeAl_3 , NiAl_3 and a quaternary compound phase were observed in equilibrium with α . FeAl_3 and the quaternary body were distinguished by etching for approximately 45 sec. in cold 10 % caustic soda solution aged before use for 24 hr. Under these conditions the quaternary phase was stained dark brown; while FeAl_3 remained comparatively unattacked. Figure 7a shows a typical ($\alpha + \text{FeAl}_3 + (\text{Fe. Co. Ni})_2\text{Al}_9$) alloy. NiAl_3 were differentiated from the quaternary by etching in a solution containing 5 % nitric acid, 3 % hydrochloric

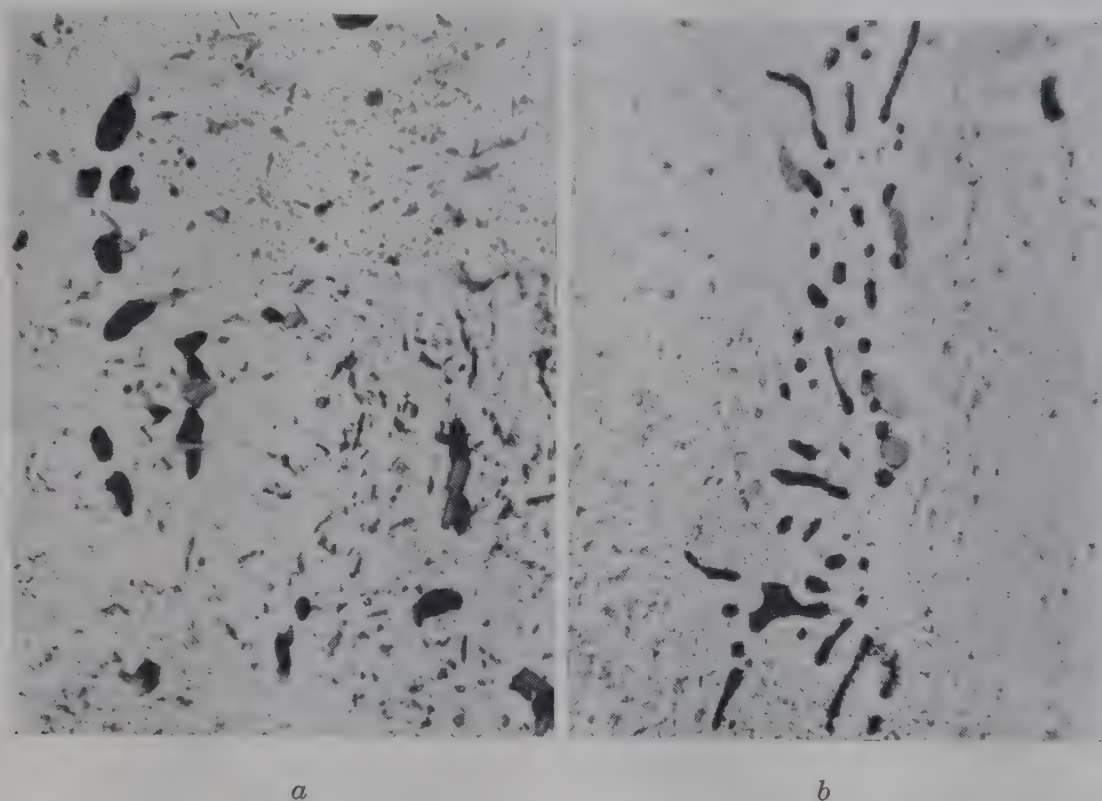


FIGURE 7. *a*, typical three-phase alloy, showing $\alpha + \text{FeAl}_3$ (light) + $(\text{Fe. Co. Ni})_2\text{Al}_9$ (dark); (magnification $\times 400$). *b*, typical three-phase alloy, showing $\alpha + \text{NiAl}_3$ (dark) + $(\text{Fe. Co. Ni})_2\text{Al}_9$ (light); (magnification $\times 400$).

acid, and 0.5 % hydrofluoric acid. Figure 7b shows a typical area of a three-phase alloy containing α (matrix), $(\text{Fe. Co. Ni})_2\text{Al}_9$ (light) and NiAl_3 (dark). All critical alloys were analyzed, and, in general, the intended compositions were closely attained. Where the analyzed composition displaced the alloys slightly from the plane of constant aluminium content at 97.5 %, correction was applied by increasing or decreasing the observed solute percentages in the same proportion to bring the corresponding points on to the appropriate plane. It will be appreciated from figure 4 that this procedure does not affect the relative positions of the plotted points in the triangular section.

Some forty-nine alloys were prepared, and annealed for 17 days at 600° C. These experiments satisfactorily established the existence of the predicted continuous ($\alpha + (\text{Fe. Co. Ni})_2\text{Al}_9$) field, and the form of the

$$(\alpha + (\text{Fe. Co. Ni})_2\text{Al}_9)/(\alpha + (\text{Fe. Co. Ni})_2\text{Al}_9 + \text{NiAl}_3)$$

boundary. The results are shown in figure 8, which is the 97.5 % aluminium section of the quaternary isothermal. At this stage the

$$(\alpha + \text{Fe.Co.Ni})_2\text{Al}_9 / (\alpha + (\text{Fe.Co.Ni})_2\text{Al}_9 + \text{FeAl}_3)$$

boundary was proved to be convex towards the Al-Fe axis; since it was of theoretical interest to define this boundary accurately, extra alloys were prepared and annealed. All critical alloys used in determining the boundary were annealed for 28 days. The

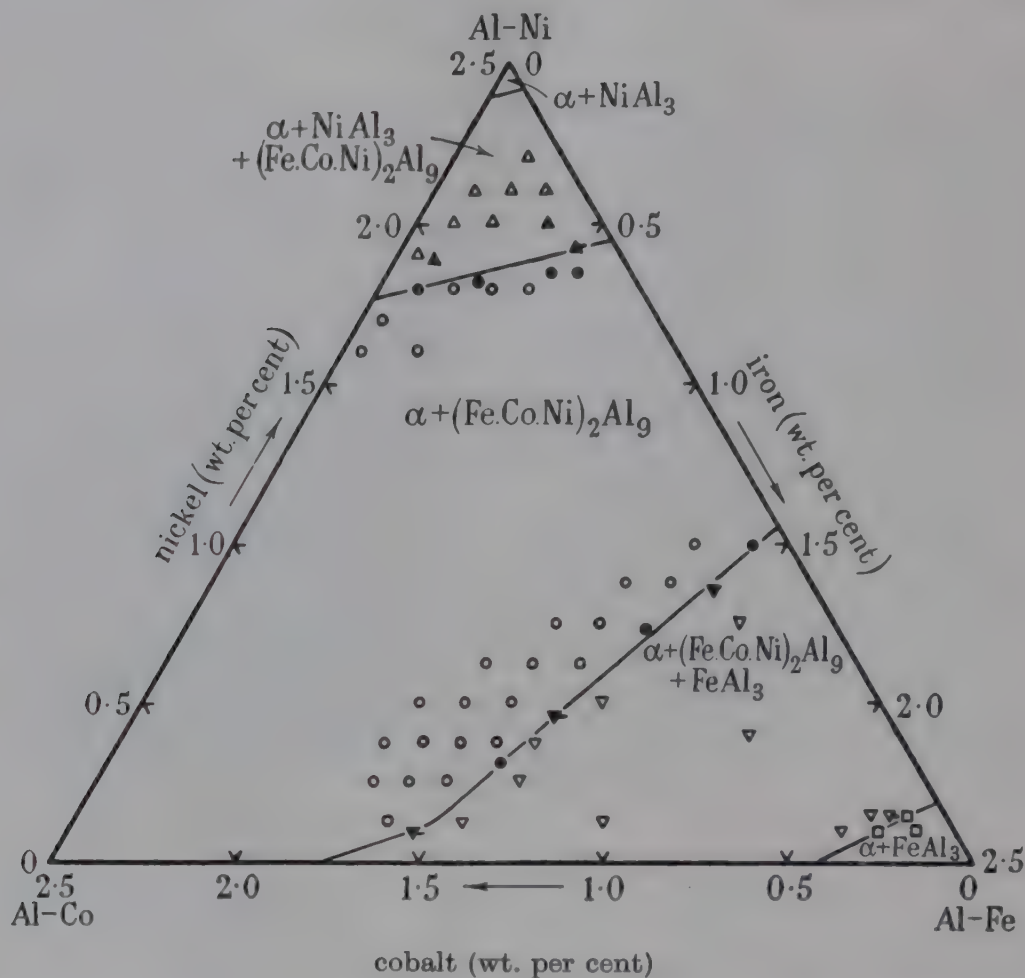


FIGURE 8. The system Al-Fe-Co-Ni; constitution at 600° C for an Al content of 97.5 %.

Δ, α + NiAl₃ + (Fe.Co.Ni)₂Al₉; ○, α + (Fe.Co.Ni)₂Al₉; ▽, α + (Fe.Co.Ni)₂Al₉ + FeAl₃; □, α + FeAl₃; ▲▼, trace of third phase; ▲●▼, analyzed composition.

results are included in figure 8, in which boundaries have been drawn in accordance with the relative amounts of constituents present in critical alloys, and show that the boundary *c'd'* consists of two intersecting straight lines. According to phase-rule principles, the boundary must be continuous, but the experimental results show that the change in direction must take place over a narrow range of compositions. The form of the boundary *c'd'* implies that the iron-rich boundary of the quaternary body itself consists, at 600° C, of two branches. The significance of this is discussed below.

The constitutions of the alloys in the (α + FeAl₃) region are consistent with a straight line joining the solubility limits of cobalt and nickel in FeAl₃ in the ternary systems.

The boundary $e'f'$ was not investigated owing to the very limited solubilities of iron and cobalt in NiAl_3 .

The equilibrium relationships in the solid state are therefore of the predicted form; the resemblance between figures 8 and 5 is immediately apparent.*

4. THE SURFACES OF PRIMARY SEPARATION

The expected forms of the quaternary surfaces of primary separation are shown in figure 6; the main feature is the continuous merging of the primary Co_2Al_9 and $(\text{Fe.Ni})_2\text{Al}_9$ fields. To examine this feature, the boundaries between the α , $(\text{Fe.Co.Ni})_2\text{Al}_9$ and FeAl_3 fields were determined, within the ranges 0 to 3.0 % cobalt, 0 to 2.5 % iron and 0 to 2.0 % nickel, by microscopical examination of slowly cooled alloys. For convenience, sections through the quaternary model of the projected fields of primary separation were prepared at 1 and 2 % nickel; the relationship of these sections to the tetrahedral model is shown in figure 9.

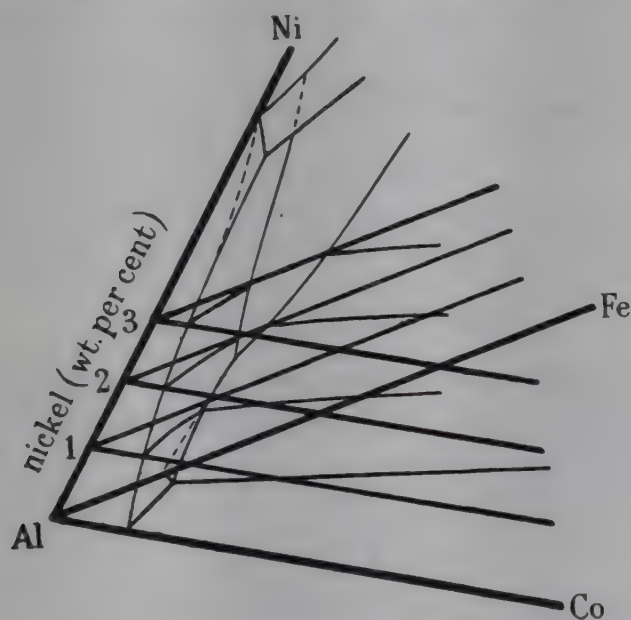


FIGURE 9. Relation of sections at constant Ni percentages to the tetrahedral model.

(i) The section at 1% nickel

Seventeen alloys, with compositions shown in figure 10, were prepared, slowly cooled, and examined; each specimen was examined before preparing subsequent alloys, in order to limit the number of experiments by deciding at each stage what compositions were most likely to give accurate information. The structures are indicated in figure 10, which also contains, for comparison, the distribution of phase fields in the ternary Al-Fe-Co system. It will be seen that the boundary between the primary α and $(\text{Fe.Co.Ni})_2\text{Al}_9$ fields has been shifted towards the aluminium-iron axis by the addition of 1 % nickel. The ternary eutectic point, involving equilibrium between liquid, α , $(\text{Fe.Co.Ni})_2\text{Al}_9$ and FeAl_3 , has moved farther away from the

* For convenience, figure 8 is plotted in weight percentages; owing, however, to the similarity of the atomic weights of iron, cobalt and nickel, transference to atomic percentages has a very small effect on the form of the phase fields.

aluminium-rich corner. This is in agreement with expectation; no intrusion of unexpected phase fields was observed.

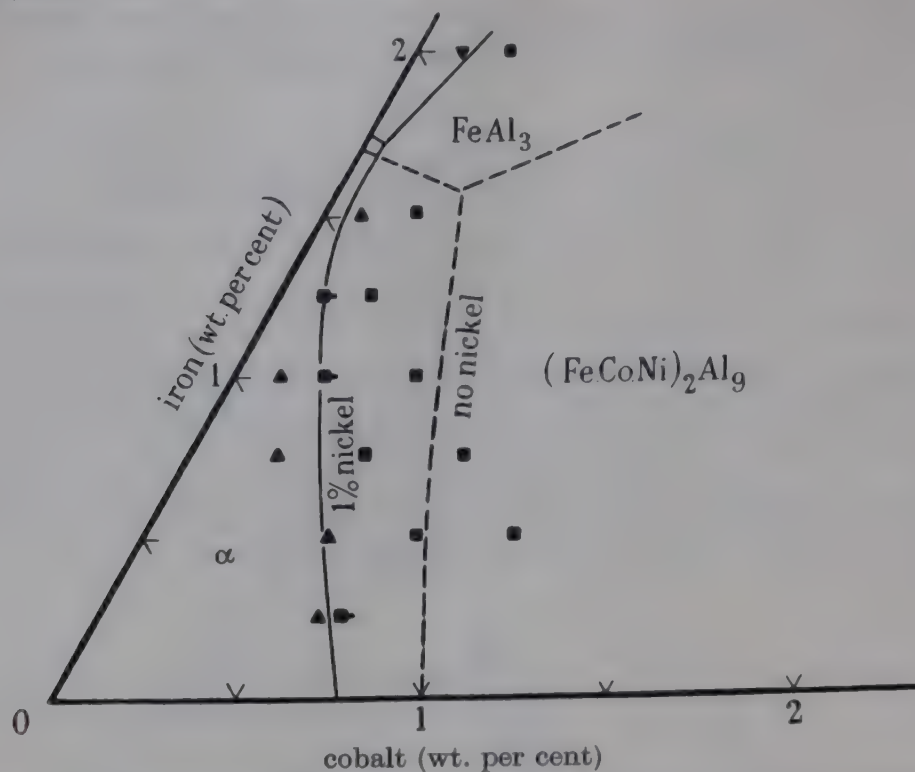


FIGURE 10. The system Al-Fe-Co-Ni; surfaces of primary separation at 1% Ni.

▼, FeAl₃; ▲, α; ■, (Fe.Co.Ni)₂Al₉; ■, trace only.

(ii) *The section at 2% nickel*

Sixteen alloys in this section were examined, and the results are shown in figure 11. At this stage the primary Co₂Al₉ field has merged smoothly with the (Fe.Ni)₂Al₉ field, in agreement with expectation.

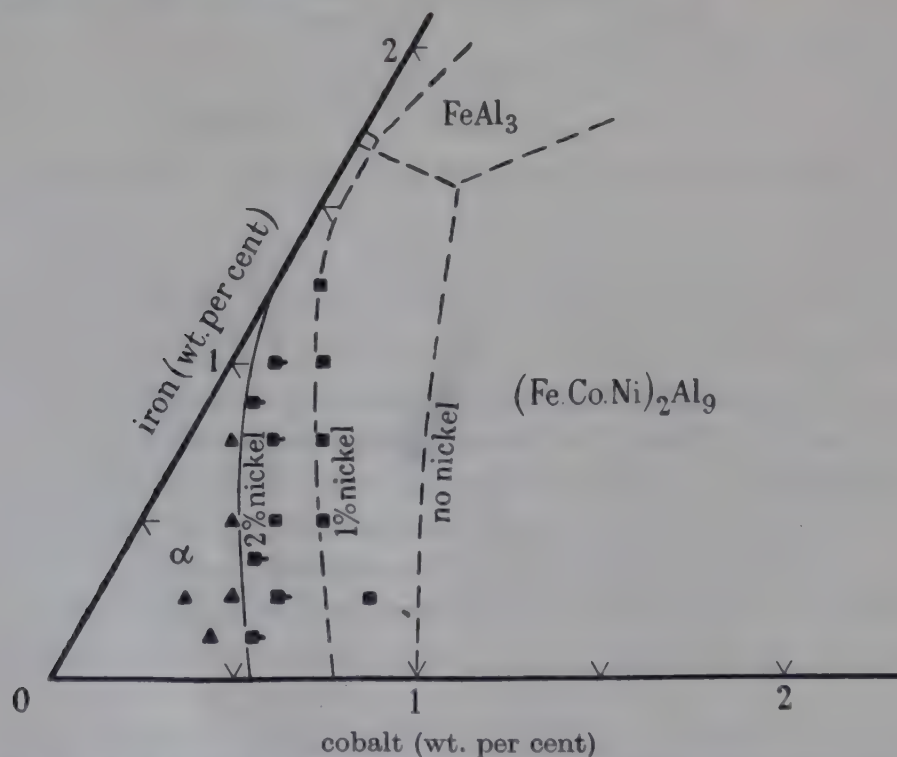


FIGURE 11. The system Al-Fe-Co-Ni; surfaces of primary separation at 2% Ni.

In both sections, the results are consistent with the intercepts derived from the previously studied ternary systems. Figure 12 shows the projections of the primary surfaces at 0, 1 and 2 % nickel superimposed; since the $(\text{Fe.Ni})_2\text{Al}_9$ and Co_2Al_9 fields have merged at 2 % nickel, the probable boundary at 3 % nickel has been derived by extrapolation, and is included, since the work described below required an approximate knowledge of its position.

The forms of the primary surfaces, therefore, are consistent with the existence of a quaternary body $(\text{Fe.Co.Ni})_2\text{Al}_9$, and enable compositions to be chosen such that this phase is obtained as primary crystals. Figure 12 was used for this purpose to obtain samples for chemical analysis and X-ray examination.

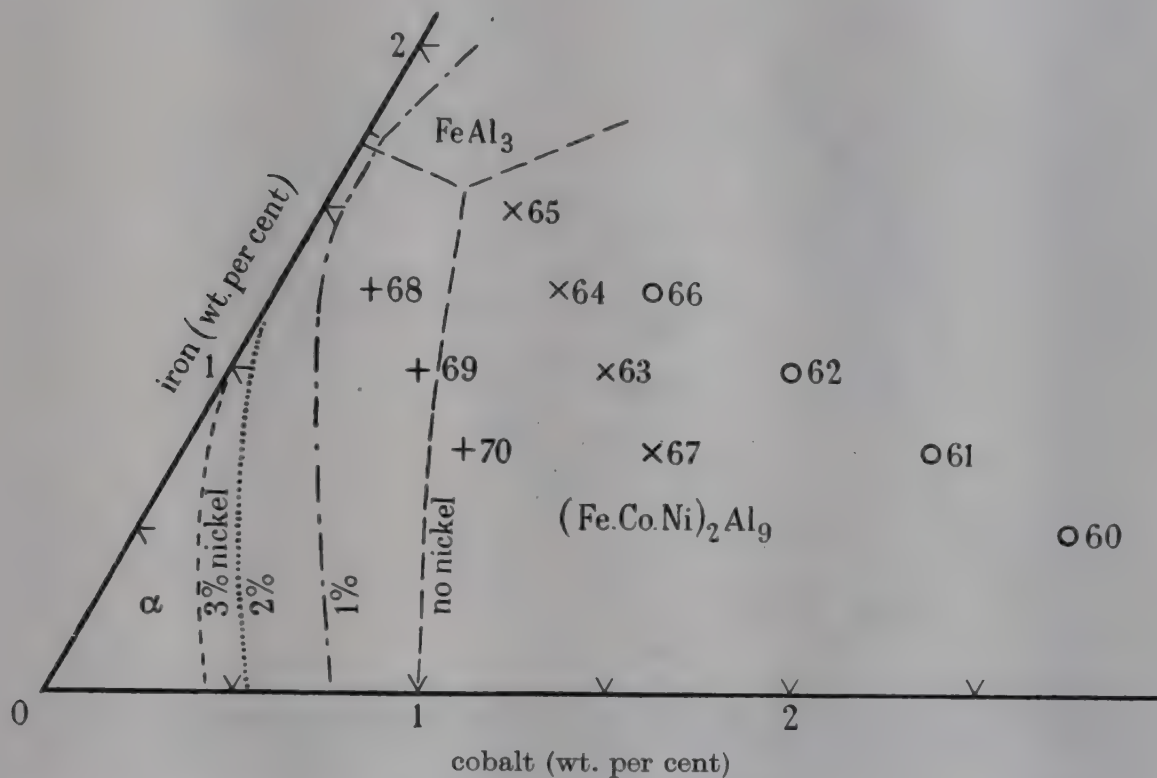


FIGURE 12. The system Al-Fe-Co-Ni; surfaces of primary separation at 1, 2 and 3 % Ni.
 — — —, nickel none; $\circ$ - - - , 1 %; $\times$ ·····, 2 %; + ·····, 3 %.

5. THE COMPOSITION OF THE QUATERNARY PHASE

In the work on the three ternary systems §2 (i), (ii), (iii) it was shown, by the extraction and analysis of primary crystals of Co_2Al_9 containing iron or nickel, and of $(\text{Fe.Ni})_2\text{Al}_9$, that the proportion of aluminium atoms to solute atoms remained constant, within the limits of experimental error, at 9:2. The X-ray diffraction patterns were also similar throughout. The same should be true of the quaternary phase; alloys containing 0 to 2.5 % cobalt, 0 to 1.5 % iron, and 0 to 4.5 % nickel, chosen on the basis of the results described in §4, were accordingly cooled slowly, and the primary crystals extracted by anodic solution of the matrix. All samples of extracted crystals were found to be $(\text{Fe.Co.Ni})_2\text{Al}_9$, thus supporting the boundaries of figure 12, and the merging of the primary Co_2Al_9 and $(\text{Fe.Ni})_2\text{Al}_9$ fields with each other.

The crystal habit of the phase varied somewhat from rectangular prisms to thin flakes of irregular acicular shape; the addition of iron favoured the formation of prisms near the aluminium-cobalt axis, while increasing proportions of nickel resulted in the occurrence of plates and thin flakes. Clean, hand-sorted samples weighing 0.2 to 0.3 g. were submitted for analysis, and the results obtained are summarized in table 1. These values are represented diagrammatically in figure 13,

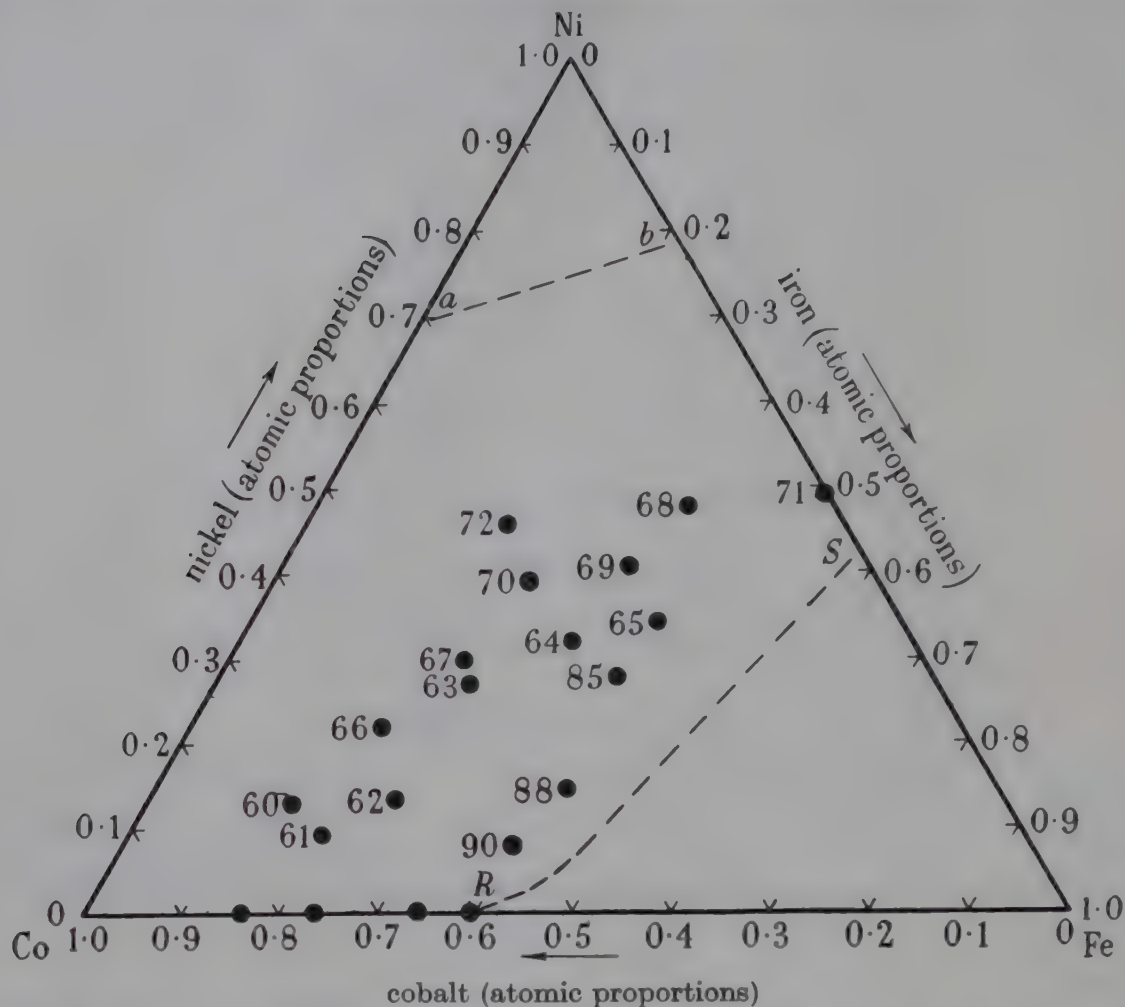


FIGURE 13. Analyzed compositions of extracted $(\text{Fe.Co.Ni})_2\text{Al}_9$ crystals; the numbers refer to table 1 and figure 12.

where they are plotted according to the atomic proportions of the solute elements; this method of plotting is unaffected by variation in the aluminium content of the samples, and is geometrically equivalent to plotting the values assuming a constant aluminium content.

The results of these experiments fully confirm the existence of a continuous solid solution field between the homogeneity ranges of $(\text{Fe.Ni})_2\text{Al}_9$ and Co_2Al_9 in the ternary systems. It will be noted from figure 13 that the region lying to the nickel-rich side of the line joining the compositions of Co_2Al_9 and FeNiAl_9 contains few points. This would appear to be owing to a tendency of the tie lines in the (liquid + primary $(\text{Fe.Co.Ni})_2\text{Al}_9$) field to run towards this line rather more than would be expected. The form of the isothermal at 600°C , however, makes it clear that the range of homogeneity on the nickel-rich side extends, at this temperature, as far as the straight line joining the homogeneity limits of Co_2Al_9 and $(\text{Fe.Ni})_2\text{Al}_9$.

In the ternary systems, a close correlation exists between the limiting compositions of the primary crystals, and the extents of the phase fields in the 600°C isothermals; a similar relationship may safely be assumed in the quaternary system, and the boundary *ab* is plotted as a straight line joining the solubility limits of nickel in extracted crystals of Co_2Al_9 and $(\text{Fe.Ni})_2\text{Al}_9$. The line *RS* is plotted in such a position as to be consistent with the previously determined solubilities of iron in primary crystals of Co_2Al_9 and $(\text{Fe.Ni})_2\text{Al}_9$ respectively.

TABLE 1

no. of alloy	type of extracted crystals	analysis (wt. %)			no. of solute atoms per 9 aluminium atoms
		Co	Fe	Ni	
60	prisms	24.88	4.83	4.36	2.15
61	needles	(25.20	6.25	3.38)	2.35 (?)
62	flakes	21.35	8.25	4.78	2.20
63	needles	16.20	8.58	9.61	2.18
64	flakes	12.22	11.36	11.44	2.24
65	flakes	8.07	13.40	11.54	2.07
66	flakes	(16.49	5.30	6.22)	1.65 (?)
67	needles	15.17	7.61	9.92	2.03
68	flakes	4.85	12.40	16.20	2.12
69	flakes	8.48	12.01	14.78	2.29
70	flakes	11.88	8.74	13.76	2.19
71	flakes	—	16.50	16.83	2.09
72	flakes	11.66	6.83	16.16	2.22
85	flakes	10.40	12.79	9.16	2.01
88	needles	13.98	12.97	4.78	1.96
90	needles	17.63	12.30	2.56	1.99

Table 1 indicates that there are small deviations from the ideal value of 2 in the ratio of solute atoms to 9 aluminium atoms. A critical examination of the results indicates, however, that there is no systematic variation of the atomic ratio with change in any one of the solute concentrations. The deviations appear to be due to statistical scatter of analytical results, which may be expected in view of the difficulties of chemical separation in the analysis of small amounts of material containing iron, cobalt and nickel present together.

The mutual replacement of iron, cobalt and nickel in the quaternary phase therefore occurs essentially atom for atom, as suggested earlier; the experiments confirm the designation of this phase as $(\text{Fe.Co.Ni})_2\text{Al}_9$, the aluminium atoms taking no part in the change of composition over the whole homogeneity range.

6. X-RAY EXAMINATION OF THE QUATERNARY PHASE

To confirm that the general crystal structure of the quaternary phase remained essentially unchanged over the whole homogeneity range, and to gain an indication of possible lattice distortion effects, a wide selection of crystals extracted for analytical experiments was examined by X-rays. Debye-Scherrer diffraction patterns were obtained, using iron $K\alpha$ radiation, from finely ground samples which had been

annealed *in vacuo* for 4 hr. at 600° C and mounted on fine glass fibres. All specimens examined gave rise to diffraction patterns characteristic of the Co_2Al_9 structure, in complete agreement with the remainder of the experiments.

In the course of the work, it was noted that although the general forms of the diffraction patterns for all specimens were very similar, the angular separations of certain lines differed slightly from alloy to alloy, giving evidence of relative changes in the dimensions of the unit cell of the monoclinic structure. This distortion of the lattice was least marked along the line joining the compositions Co_2Al_9 and FeNiAl_9 , and most marked with nickel or iron in excess of the proportion 1:1. A separate study of these effects is being made.

7. DISCUSSION

The results of the experimental work show that the theoretical predictions concerning the constitution of the aluminium-rich region of the quaternary Al-Fe-Co-Ni system are satisfactorily confirmed. In particular, it is established that:

(a) The process of substituting cobalt for iron and nickel in $(\text{Fe.Ni})_2\text{Al}_9$ continues without interruption until the composition Co_2Al_9 , on the binary Al-Co axis, is reached. Thus there is a quaternary phase $(\text{Fe.Co.Ni})_2\text{Al}_9$ stable over a wide homogeneity range within which the transitional metals are interchangeable while the aluminium content remains constant. There are small changes of lattice spacing accompanying the composition changes.

(b) The surfaces of primary Co_2Al_9 separation in the Al-Co-Fe and Al-Co-Ni diagrams merge simply into the primary $(\text{Fe.Ni})_2\text{Al}_9$ field in the Al-Fe-Ni diagram.

(c) The constitution at 600° C is of the expected form; there is a continuous $(\alpha + (\text{Fe.Co.Ni})_2\text{Al}_9)$ field joining the homogeneity ranges of Co_2Al_9 and $(\text{Fe.Ni})_2\text{Al}_9$ in the component ternary systems. The other equilibria agree with expectation, and the boundaries of the $(\alpha + (\text{Fe.Co.Ni})_2\text{Al}_9)$ region give interesting information concerning the stability of the quaternary body.

(d) No new phases enter into equilibrium with the aluminium-rich primary solid solution.

It has therefore proved feasible to predict the equilibrium relations in the quaternary system from those in the three component ternary systems, two of which were themselves largely predicted as a result of the original theoretical interpretation of the Al-Fe-Ni alloys, §2 (i), (ii), (iii).

Figure 8, which shows the constitution at 600° C for an aluminium content of 97.5 %, is of further interest. Owing to the negligible solid solubilities of iron, cobalt and nickel in aluminium, the form of the $(\alpha + (\text{Fe.Co.Ni})_2\text{Al}_9)$ field corresponds geometrically with the homogeneity range of the quaternary phase itself at 600° C. It follows that the nickel-rich boundary of the $(\text{Fe.Co.Ni})_2\text{Al}_9$ phase is a straight line; along this line the electron:atom ratio, as expected, is close to that at which the Co_2Al_9 structure itself, on adding further electrons, becomes unstable relative to possible alternatives. The boundary representing the limiting solubility of iron in the quaternary phase at 600° C must have the form of the

$$(\alpha + (\text{Fe.Co.Ni})_2\text{Al}_9)/(\alpha + (\text{Fe.Co.Ni})_2\text{Al}_9 + \text{FeAl}_3)$$

boundary in figure 8. Co_2Al_9 itself, of electron:atom ratio 2.125, dissolves iron (at 600°C) until the electron:atom ratio is reduced to 2.07; FeNiAl_9 , however, in which half the cobalt sites are occupied by iron and half by nickel, can only support a reduction of the electron:atom ratio from 2.16 to 2.14. Figure 8 shows that the iron-rich boundary of the quaternary phase at 600°C is not a straight line joining the iron-rich limits of $(\text{Fe.Ni})_2\text{Al}_9$ and Co_2Al_9 . The corresponding boundary of the $(\alpha + (\text{Fe.Co.Ni})_2\text{Al}_9)$ field moves away from the Al-Fe-Co face of the model in a direction which is consistent with the quaternary phase maintaining the electron:atom of 2.07 along the iron-rich limit of its homogeneity range. After a short distance, however, the direction of the boundary changes; the electron:atom ratio rises uniformly until the limiting composition of the $(\text{Fe.Ni})_2\text{Al}_9$ phase is reached at the value 2.14. The composition at which the boundary departs from the constant electron:atom ratio of 2.07 is that at which exactly half the cobalt sites in the Co_2Al_9 structure are replaced by iron and nickel. Up to this point, therefore, the quaternary phase remains as stable towards a reduction of the number of electrons per atom as the binary phase Co_2Al_9 . With greater amounts of iron and nickel, the structure becomes progressively less able to maintain a deficit of electrons.

In the ternary Al-Fe-Co system, Co_2Al_9 enters into a eutectic reaction with the compound FeAl_3 ; in the Al-Fe-Ni system, however, $(\text{Fe.Ni})_2\text{Al}_9$ is formed peritectically from FeAl_3 . This has been taken as an indication that $(\text{Fe.Ni})_2\text{Al}_9$ is less stable than Co_2Al_9 , and therefore less able to withstand an electron deficit. It is of interest to note in confirmation that the composition of the quaternary phase which separates from the liquid at the point at which the peritectic relationship changes over into the eutectic relationship also corresponds with the replacement of half the cobalt sites by iron and nickel.* The transition from the eutectic relationship to the peritectic relationship, and the deviation of the iron-rich boundary of the quaternary phase from a constant electron:atom ratio, are both associated with the replacement of half the cobalt sites in the structure.

The experiments described in the present paper confirm that very close relationships exist between the quaternary system Al-Fe-Co-Ni and the three component ternary systems. In particular, the details of the nickel-rich and iron-rich boundaries of the range of homogeneity of the quaternary phase confirm the importance of the electron:atom ratio in this type of intermetallic phase, the electron:atom ratios being calculated on the basis of the acceptance of electrons by the transitional metal atoms. The work also indicates that, although the electron:atom ratio is an important factor, other factors influence the stability of phases of the type of Co_2Al_9 . The fact that both Co_2Al_9 and $(\text{Fe.Ni})_2\text{Al}_9$ dissolve nickel to the same electron:atom ratio emphasizes the electronic aspect, and strongly suggests that such compounds may legitimately be considered as analogous to electron compounds, characterized by Brillouin zones which can accommodate at relatively low energies the number of electrons present. The behaviour on reduction of the electron:atom ratio emphasizes the structural aspect. Since the atomic diameters of iron, cobalt and nickel are similar, it is unlikely that simple lattice distortion effects will account for the relatively sudden change in

* Details of the work carried out to examine the transition from the peritectic to the eutectic relationship are to be published elsewhere.

the ability of the quaternary structure to withstand an electron deficit. It is more probable that selective replacement of cobalt atom sites occurs up to the point at which half these sites are occupied by other transitional metal atoms. Further replacements may then result in a more random replacement with a consequent weakening of the structure.

This research was carried out in the Metallurgy Department of the University of Birmingham, under the general supervision of Professor D. Hanson, D.Sc., to whom the authors' thanks are due for his interest and support. The authors must also acknowledge the valuable assistance of Mr G. Welsh in the experimental work. Analyses of the annealed alloys and the extracted crystals were carried out by Messrs Johnson Matthey and Co. Ltd., or by Mr J. Hawkes and Mr P. Wingrove of the authors' Department.

The authors express their gratitude to the Department of Scientific and Industrial Research, the Royal Society, the Chemical Society, and Imperial Chemical Industries Ltd., for general financial assistance.

REFERENCES

- Bradley, A. J. & Taylor, A. 1940 *J. Inst. Metals*, **66**, 53.
Douglas, A. M. B. 1948 *Nature*, **162**, 566.
Pauling, L. 1938 *Phys. Rev.* (ii), **54**, 899.
Phillips, H. W. L. 1942 *J. Inst. Metals*, **68**, 34.
Raynor, G. V. & Pfeil, P. C. L. 1946-7a *J. Inst. Metals*, **73**, 397.
Raynor, G. V. & Pfeil, P. C. L. 1946-7b *J. Inst. Metals*, **73**, 609.
Raynor, G. V. & Waldron, M. B. 1948 *Proc. Roy. Soc. A*, **194**, 362.
Raynor, G. V. & Waldron, M. B. 1949 *Phil. Mag.* [vii], **40**, 198.
Schrader, A. & Hanemann, H. 1943 *Aluminium*, **25**, 339.

Identification of microseismic activity with sea waves

BY J. DARBYSHIRE, *Royal Naval Scientific Service*

(Communicated by G. E. R. Deacon, F.R.S.—Received 9 March 1950)

Three series of simultaneous wave and microseism records are examined. They give a clear indication that bands of microseismic waves from different sources can be distinguished by submitting seismograph records to frequency analysis. The agreement between the results of analysis and the theoretical expectation from the prevailing meteorological conditions appears to justify the assumption that microseismic waves of different periods travel independently. Under the simple meteorological conditions that have been studied, each band of microseismic activity can be identified with a band of sea waves of twice its period. The existence of this two to one ratio between the period of waves and microseisms affords some confirmation of the theory that microseisms are produced in a region of interference between similar wave trains travelling in opposite directions either near the coast or in deep water.

SEA WAVES AND MICROSEISMS

Bernard (1941) obtained evidence which led him to believe that microseismic oscillations have periods which are half those of the sea waves which give rise to them. He reached this conclusion by comparing the microseisms recorded at Averroes near Casablanca with simultaneous observations on the waves reaching the coast. The same ratio was noticed by Deacon (1947) between microseisms recorded at Kew and sea waves recorded at Perranporth on the north coast of Cornwall. The reason for an exact relationship became apparent when Longuet-Higgins & Ursell (1948), using theoretical work on standing waves by Miche (1944), showed that the mean pressure on the sea bottom below a region of interference between trains of waves of the same period travelling in opposite directions varies with half the wave period. The mean pressure variation is a second-order effect proportional to the product of the amplitudes of the two wave trains and the square of the frequency, and unlike the pressure variation below a single train of waves, it does not vanish with depth. This result has been used by Longuet-Higgins (1950) as the basis of a new theory of microseism generation. He shows that interference between trains of waves of the same dominant period, travelling in opposite directions, causes variations in mean pressure on the sea bottom sufficiently large to account for the observed microseismic oscillations, even when the water is very deep. Such interference can occur when waves or swell are reflected from a steep coast, when waves generated in different quadrants of a fast-moving depression meet each other travelling in opposite directions, and when similar trains of swell travel in opposite directions from two storm areas.

There may, according to the theory, be some departure from the exact two to one ratio between wave and microseism periods when each wave band is of appreciable width. For a given wave height, more energy will be communicated to the sea bottom by short waves than by long waves. Resonance of the overlying water column may also favour a frequency which corresponds to a compression wave in water which is $(\frac{1}{2}n + \frac{1}{4})$ times as long as the depth, n being an integer.

Comparison of the variation in microseismic activity at Kew with changes in wave height at Perranporth (Deacon 1947) gave a clear indication that waves entering the coastal region west of the British Isles are mainly responsible for the microseismic activity at Kew, but similar work in the U.S.A. (Ramirez and Gilmore) showed that the microseismic activity recorded in observatories in the eastern states was due mainly to the presence of storms over deep water. The two generalizations can be reconciled if it is admitted that the microseismic activity at a particular observatory is due to more than one source. In the British Isles the microseismic activity is not due entirely to wave interference in the coastal region but also to interference of waves in distant regions. If the wave periods in the coastal region are appreciably different from those in the distant region, there should be two bands of microseism frequencies. The method of investigation previously used, the visual examination of records, is not sufficiently precise to distinguish weaker oscillations from one source in the presence of stronger oscillations from another source; but it should be possible to separate them by using the technique of frequency analysis. This method was applied by Barber & Ursell (1948) to sea waves. They found sufficiently close agreement between their results and those predicted from the hydrodynamical theory to justify their assumption that wave trains of different periods travel independently, the velocity potentials being additive.* In this paper, it is proposed to make a similar assumption that microseismic waves travel independently, and its validity will be tested by the agreement which is found between theory and the results of analysis.

The microseism records were converted by the use of a photo-electric curve-following machine, described by Tucker & Collins (1947, 1949), into a form suitable for analysis on the machine described by Barber, Ursell, Darbyshire & Tucker (1946). This involved the magnification of each microseism record about twenty-five times. This magnification may introduce spurious fluctuations in the final record which may be partly responsible for the background of Fourier peaks at the long-period end of the periodograms, but examination of analyses of records taken at successive times shows that within the frequency range covered by the microseisms, the envelope of the microseismic waves stands out sufficiently above the background to give a reasonable indication of the main frequencies present. The height of the envelope affords a measure of the amplitude of the microseisms; the energy in any period interval will be proportional to the sum of the squares of the peaks on the periodogram within that interval.

Three examples are presented in which a single depression over the Atlantic Ocean was producing large waves at a time when the wave activity in the coastal region was small and of different period. The seismograph records were routine recordings taken at Kew Observatory with the Galitzin vertical seismograph. Simultaneous wave records were obtained from Perranporth for the first two examples and from Pendeen for the third. The waves were measured by pressure recorders laid on the sea bed, the instrument at Perranporth lying at a depth of 45 ft. and that at Pendeen at 100 ft. The frequency-response curve of the Kew Galitzin vertical seismograph is nearly uniform

* The method has been criticized by Seiwel & Wadsworth (1949) and Seiwel (1949) but their criticisms are answered by Barber and Ursell (1948).

between 6 and 13 sec. In comparing the mean periods of the bands of waves and microseisms, some account must be taken of the effect of tidal streams on the wave period as measured by a stationary instrument. Barber & Ursell (1948) and Barber (1949) have shown that the recorded period can vary by as much as 1 sec. on either side of the true period because the tidal stream changes direction relative to the

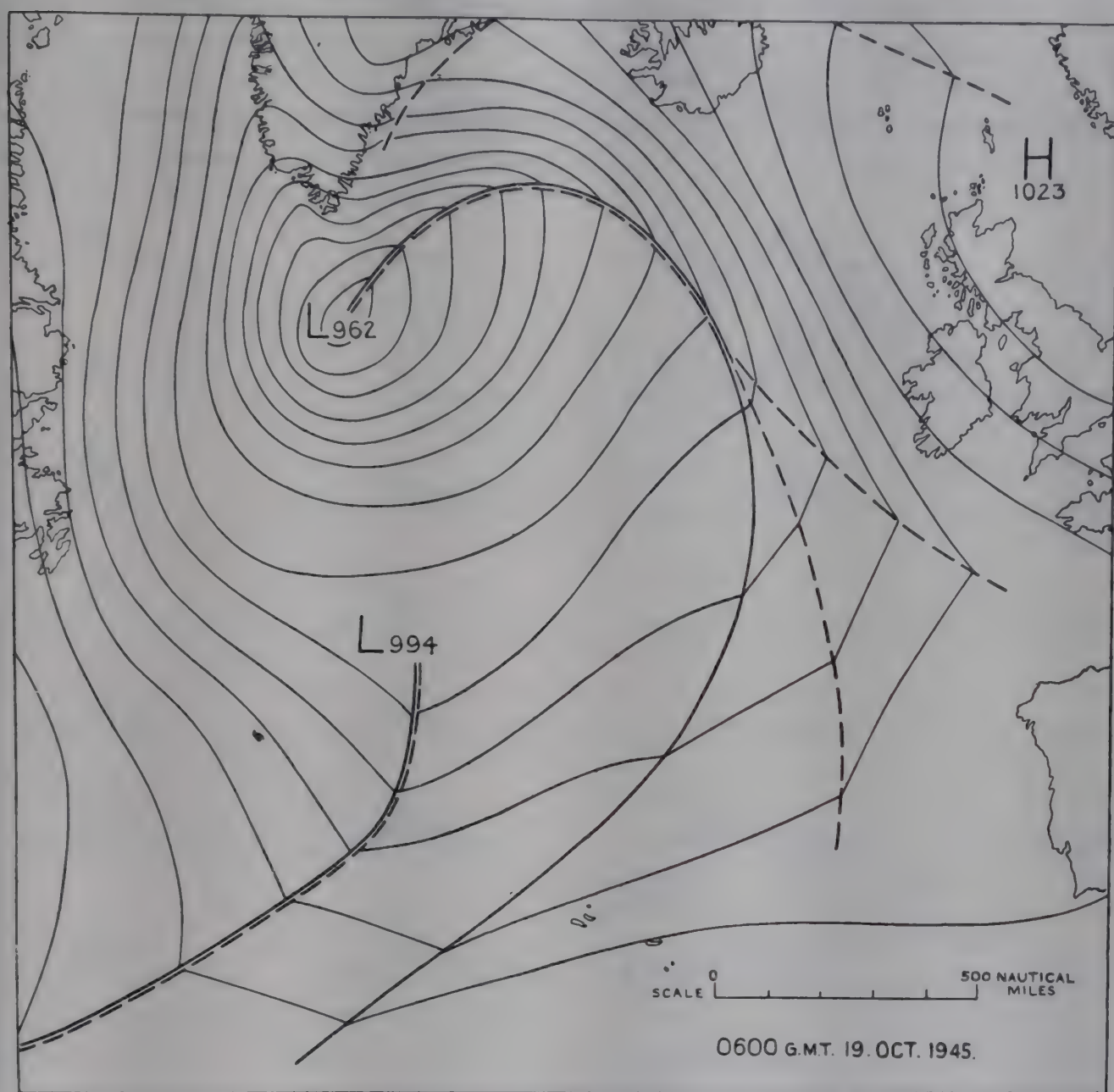


FIGURE 1. Meteorological chart of the North Atlantic Ocean, 19 October 1945.

direction of wave propagation. The change of stream gives rise to a semi-diurnal oscillation of approximately 1 sec. amplitude superimposed on the general trend of a wave band from long to short periods. The wave band is usually wide enough to include waves of the same period among the incident and reflected waves, and it is presumably these which produce the microseisms; if so, the periods of the microseisms should not be affected by variations in tidal streams, and it is to be expected that the trend of the upper and lower limits of the wave band from record to record should therefore be more regular for microseisms than for waves.

18 to 22 October 1945

In this example, there is a deep depression approximately 300 miles south of Greenland. The meteorological chart for 06.00 hr. 19 October is reproduced in figure 1, and frequency analyses of approximately 20 min. records of waves and microseisms

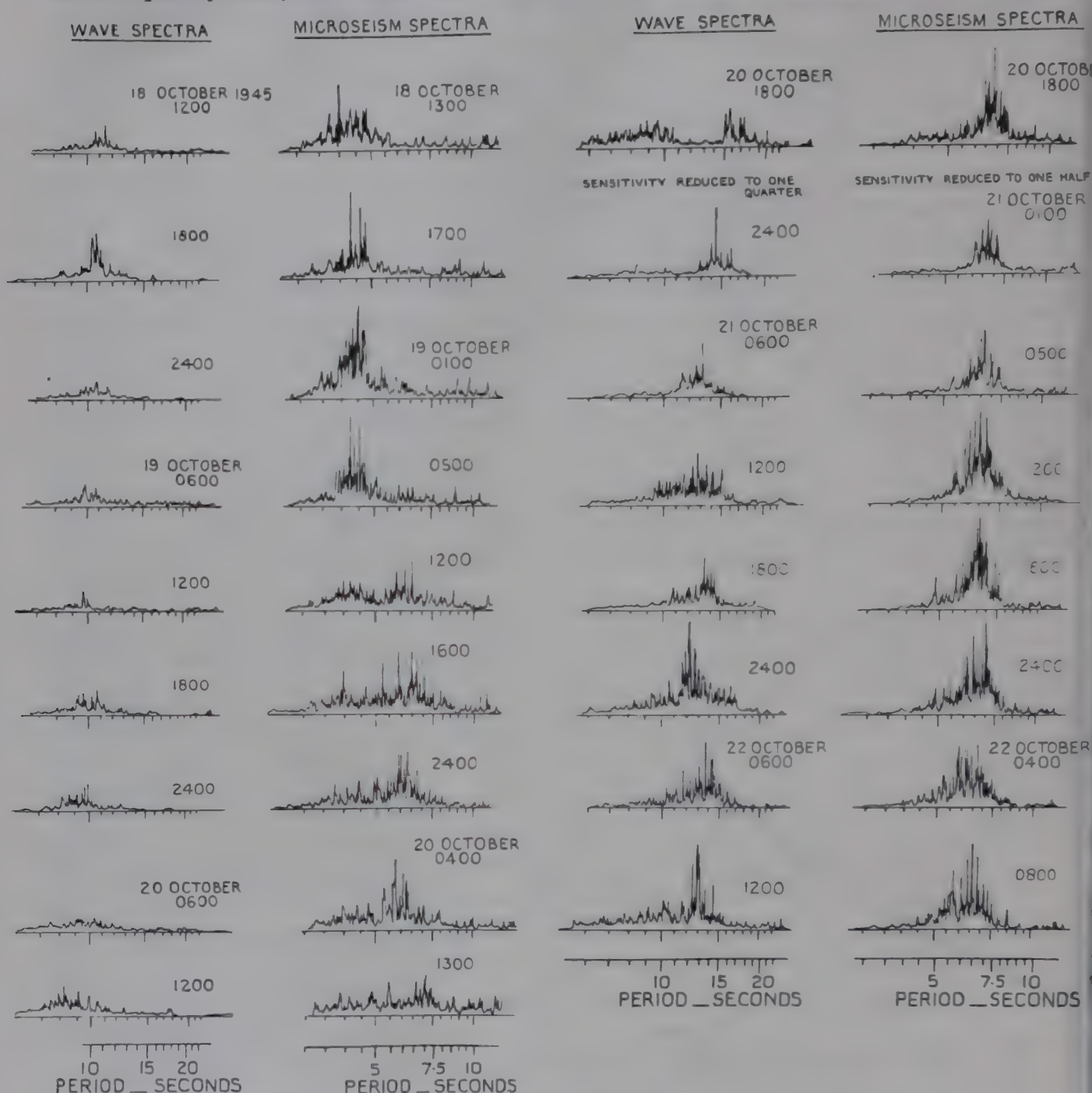


FIGURE 2. Simultaneous wave and microseism spectra, 18 to 22 October 1945.

are shown in figure 2. The first four pairs of analyses show the presence of one band of wave periods and one band of microseism periods with mean periods of 10 to 11 and 4 to 5 sec. respectively. A second band of microseism activity with a mean period of 6 to 6.5 sec. began to appear at 12.00 hr. 19 October, and the absence of a corresponding band in the wave analyses suggests that the new microseism activity had its origin outside the coastal region west of the British Isles. The storm south of Greenland was the only prominent wave-generating area, and the 12 to 13 sec. waves

required by the new theory could only be generated there. The wave interference required by the theory was probably occurring where the waves generated in the northern sector of the storm were reflected from the east coast of Greenland. Swell from the southern sector of the storm began to arrive at Perranporth between 12.00 and 18.00 hr. 20 October, after an interval which is approximately the time required for the travel of the first waves generated in the storm, and as soon as it reached the coastal region, a third band of microseismic activity developed with a mean period of 6 to 8 sec. This could be attributed to interference in the coastal region between incoming and reflected waves. The mean period of the swell when it first arrived was 15 to 18 sec., but it soon decreased to between 12 and 16 sec. in good agreement with the two to one ratio.

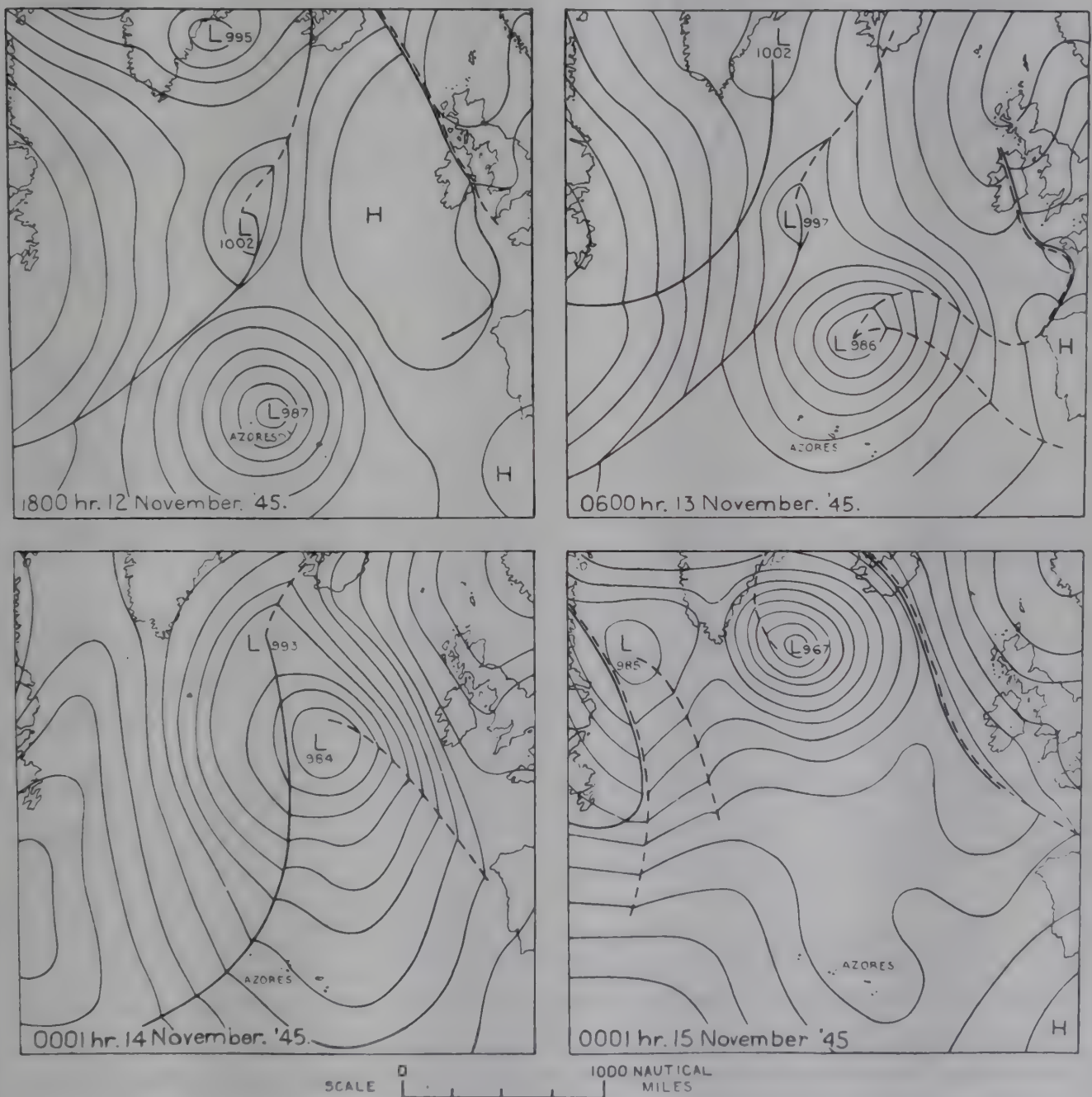


FIGURE 3. Meteorological charts of the North Atlantic Ocean, 12 to 15 November 1945.

12 to 16 November 1945

This example was chosen to study the effect of a fast-moving depression over deep water. The depression developed west of the Azores on 12 November and moved along a curved path north-east, north and north-west as indicated by the four charts

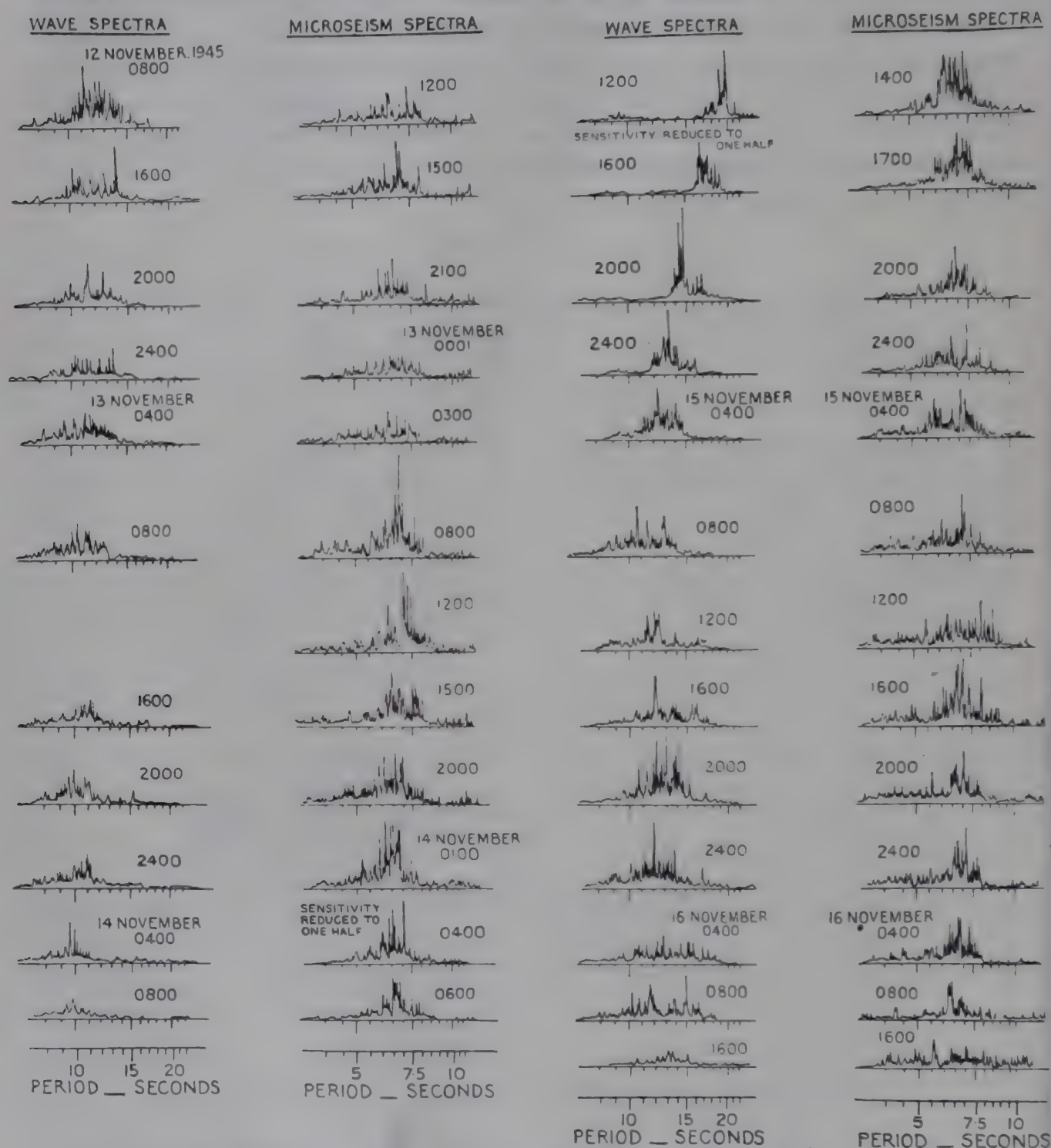


FIGURE 4. Simultaneous wave and microseism spectra, 12 to 16 November 1945.

in figure 3. Corresponding analyses of 20 min. wave and microseism records are shown in figure 4. The first five analyses show one wave band of period 10 to 14 sec., and one microseism band of period 6 to 8 sec. Between 03.00 and 08.00 hr. 13 November, there was a marked increase in the microseisms without any corresponding increase in the wave activity in the coastal region. Except for a moderate

sea in the Davis Strait, the storm near the Azores was the only active wave-generating area, and conditions likely to produce interference between two trains of waves of 12 to 16 sec. period travelling in opposite directions were associated with it. Examination of the charts in figure 3 shows that there was a large sea area north of the Azores in which the wind backed from east to west between 18.00 and 06.00 hr. 12 to 13 November, and it appears probable that interference between similar wave trains generated in the outer parts of the area would take place at about the same time as the increase in microseismic activity shown in figure 4. Similar wave interference could be expected along the subsequent path of the depression, and the further increase in microseismic activity after 01.00 hr. 14 November could be attributed to the widening of the region of wave interference and its closer approach to the British Isles. The swell from the storm began to arrive between 08.00 and 12.00 hr. 14 November, and there was an increase in microseismic activity during this period. The mean period of the microseisms after the increase at 08.00 hr. 13 November, 6 to 7.5 sec., was very nearly half that of the waves of 11 to 14 sec. period which arrived at Perranporth 1 day later, the approximate travel time from the probable region of wave interference north of the Azores. The first swell to arrive had a mean period of 18 to 20 sec., but there were no microseismic waves of 9 to 10 sec. period to correspond with it. A possible explanation is that the effect of such long low swell would be small compared with that of the heavier swell of shorter period. Except for this onset of the swell, the wave and microseism periods are in the ratio of two to one. The example shows that the analyses allow an individual band of microseismic waves, or a change of activity, to be identified with the wave activity in the coastal regions and the Atlantic Ocean.

12 to 16 March 1945

This example is concerned with a deep depression which developed in approximately 42° N, 45° W in the early morning of 12 March and reached a position 55° N, 40° W 2 days later. Meteorological charts for these two times are shown in figure 5, and the wave and microseism spectra for 12 to 16 March in figure 6. Because of the great magnification required in this case to bring the seismological records into an analyzable form, the background on the analyses is greater than usual, but there is no doubt about the presence during the first 2 days of a band of microseisms of mean period 5 to 7 sec., which is half the period of the swell reaching Pendeen. A second band of microseisms with a mean period of approximately 9 sec. began to arrive at 23.00 hr. 13 March, and this fresh activity can only be due to the storm shown in figure 5. There is, however, some uncertainty about the precise area in which the microseisms are generated.

Low swell with a period as much as 23 sec. began to arrive at Pendeen at 13.00 hr. 14 March, and much higher swell of approximately twice the period of the new microseisms followed at about 23.00 hr. If the microseismic waves which appeared 24 hr. earlier had been caused by this swell being involved in wave interference on its way to Pendeen, the region of interference could not be more than 720 miles from Pendeen. The meteorological conditions give no reason to suppose that two trains of waves of such long period could have been travelling in opposite directions so near

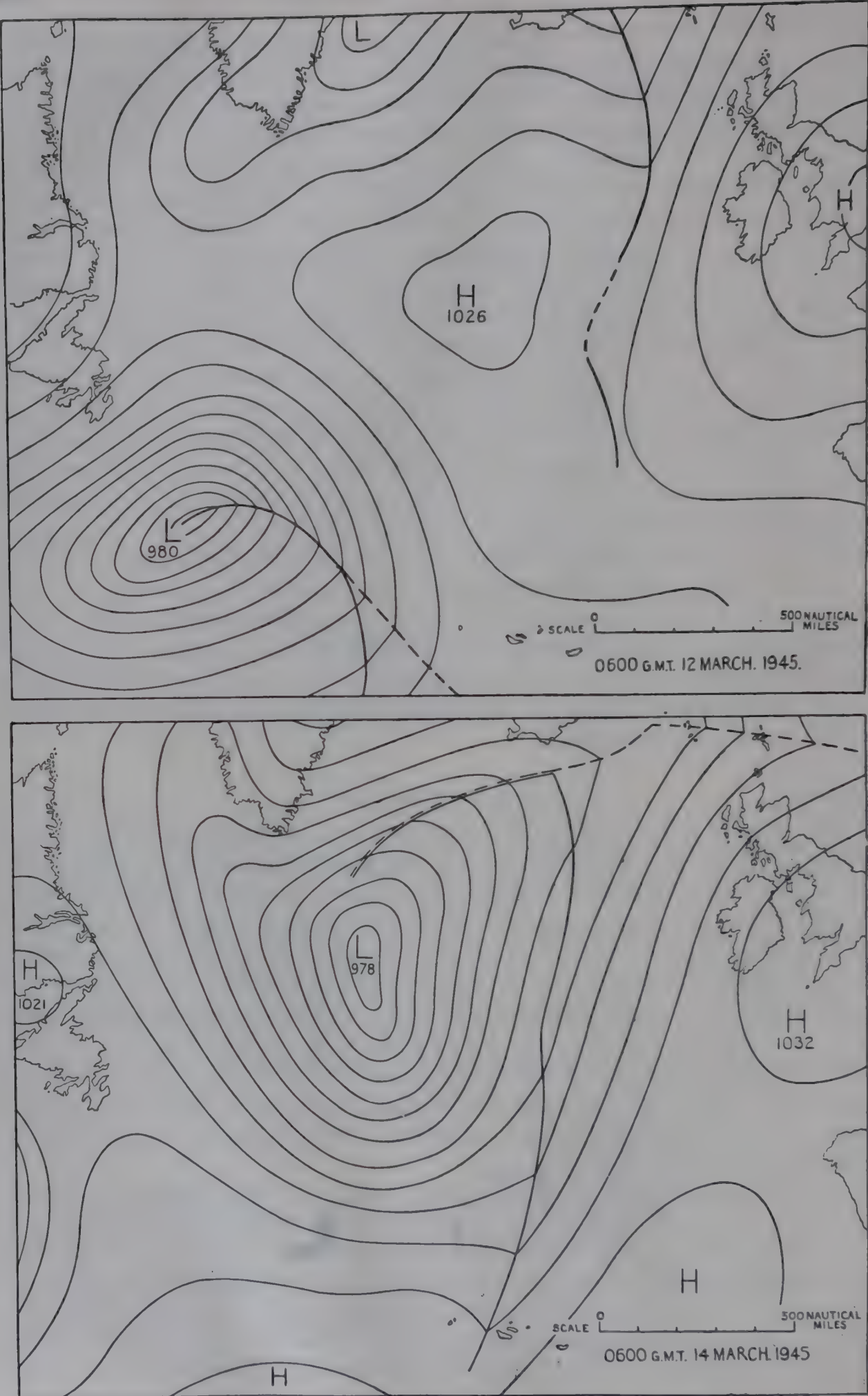


FIGURE 5. Meteorological charts of the North Atlantic Ocean, 12 and 14 March 1945.

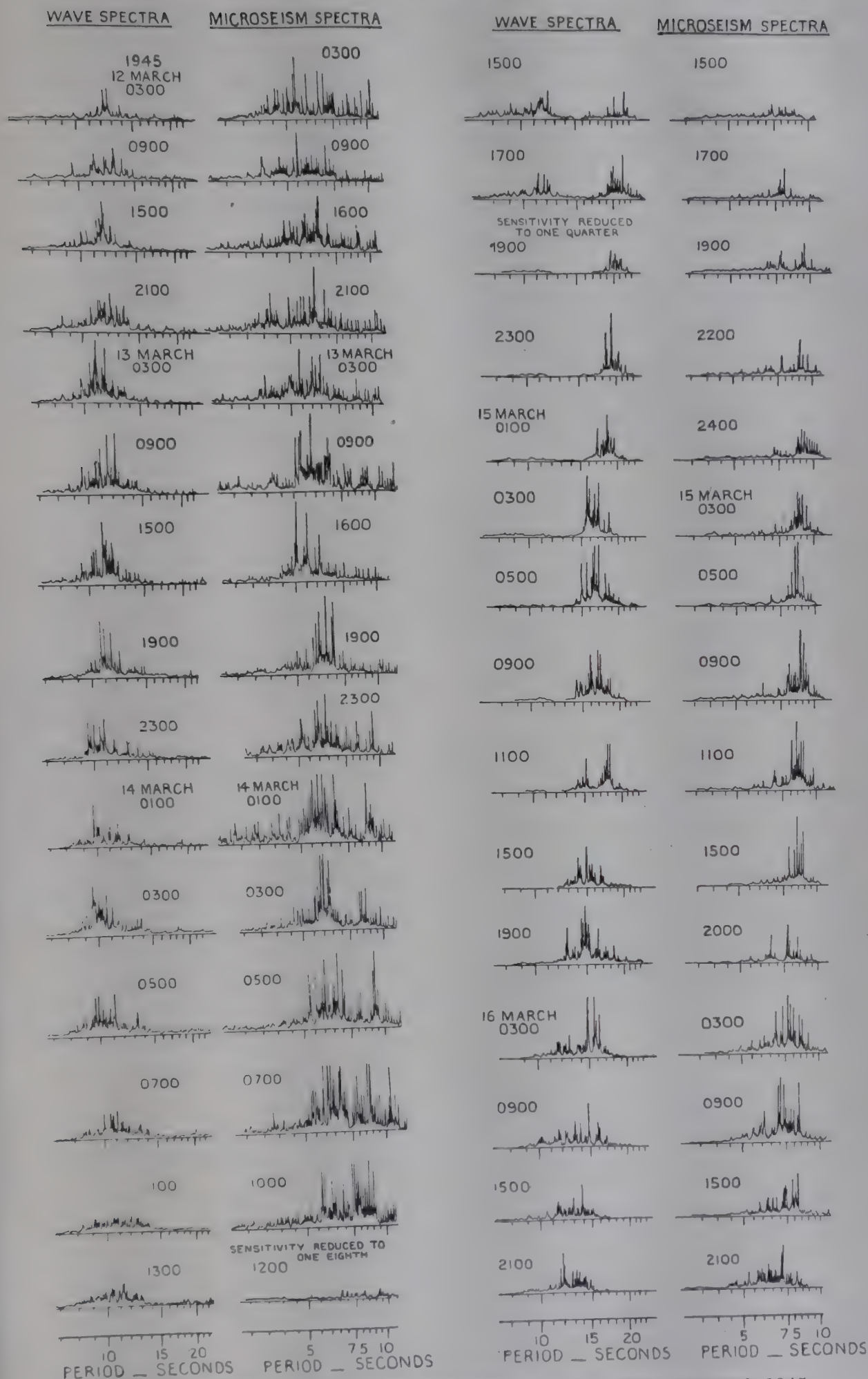


FIGURE 6. Simultaneous wave and microseism spectra, 12 to 16 March 1945.

the British Isles, and it is more reasonable to conclude that the microseisms were generated when the swell from the northern sector of the storm, as shown in figure 5, was being reflected from the east coast of Greenland. A wave prediction based on this and the preceding charts indicated that waves of 18 to 20 sec. period would reach the coast of Greenland at approximately 18.00 hr. 13 March, and their height would be about 15 ft.

There was a further increase in microseism activity when the swell from the storm reached the coast of Cornwall. The longest microseisms had a mean period of 9 to 10 sec., and there is no clear indication that the first low swell of more than 20 sec. period produced longer oscillations. The narrowness of the wave and microseism bands on 15 March affords a remarkable opportunity of checking the two to one ratio between the dominant periods, and, within the limits of accuracy set by the measurements and the tidal streams, the theoretical ratio is well upheld. The microseism band broadens towards shorter periods like the waves, and the mean periods show the same trend.

CONCLUSION

The three series of simultaneous wave and microseism records which have been presented give a clear indication that bands of microseismic waves from different sources can be distinguished by submitting seismograph records to frequency analysis. The agreement between the results of analysis and the theoretical expectation from the prevailing meteorological conditions appears to justify the assumption that microseismic waves of different periods travel independently. Under simple meteorological conditions each band of microseismic activity can be identified with a band of sea waves of twice its period. The existence of this two to one ratio between the period of waves and microseisms affords some confirmation of the theory that microseisms are produced in a region of interference between similar sea wave trains travelling in opposite directions either near the coast or in deep water.

The author expresses his gratitude to the Superintendent of Kew Observatory for the loan of the seismograph records on which the investigation is based; to Dr Beer, formerly of the staff of Kew Observatory; to Dr G. E. R. Deacon and other members of the Oceanographical Group in the Royal Naval Scientific Service for advice and assistance in preparing the paper; and to the Admiralty for permission to publish the work.

REFERENCES

- Barber, N. F. 1949 *Proc. Roy. Soc. A*, **198**, 81.
 Barber, N. F. & Ursell, F. 1948 *Phil. Trans. A*, **240**, 527.
 Barber, N. F., Ursell, F., Darbyshire, J. & Tucker, M. J. 1946 *Nature*, **158**, 329.
 Bernard, P. 1941 *Bull. Inst. océanogr. Monaco*, **38**, 800.
 Deacon, G. E. R. 1947 *Nature*, **160**, 419.
 Gilmore, M. H. 1946 *Bull. Seism. Soc. Amer.* **36**, 89.
 Longuet-Higgins, M. S. 1950 *Phil. Trans.* (in the Press).
 Longuet-Higgins, M. S. & Ursell, F. 1948 *Nature*, **162**, 700.
 Miche, M. 1944 *Ann. Ponts Chauss.* **2**, 42.
 Ramirez, J. E. 1940 *Bull. Seism. Soc. Amer.* **30**, 35, 137.
 Seiwel, H. R. 1949 *Proc. Nat. Acad. Sci., Wash.*, **35**, 518.
 Seiwel, H. R. & Wadsworth, G. P. 1949 *Science*, **109**, 271.
 Tucker, M. J. & Collins, G. 1947 *Electronic Engng*, **19**, 398.
 Tucker, M. J. & Collins, G. 1949 *Ass. Séis. Météorol. Océanog. Phys., Séance Communc*, p. 16.

Constitution and mechanism of the selenium rectifier photocell

By J. S. PRESTON, *Light Division, National Physical Laboratory*

(Communicated by Sir Charles Darwin, F.R.S.—Received 3 December 1949—

Revised 28 March 1950)

This paper describes experiments to elucidate the exact physical and chemical structure of the selenium rectifier photocell, especially that of the thin surface film. A technique is described for sputtering films of cadmium oxide which, though transparent in the thickness required for a cell, have an electrical conductivity exceeding that of graphite. The thickness of the films can be closely controlled. With such films, on pure crystalline selenium, cells were produced with white-light sensitivities of over 700 μA per lumen, open-circuit voltages up to 0.5 V under high illumination, and maximum quantum efficiencies up to 70 %.

The optical properties of the films are described, and the way in which the technique may be used to produce other non-metallic films is indicated. The cadmium oxide is found to have a negative Hall coefficient, and is therefore an *N* type semi-conductor.

Further experiments with single films of gold, and double films of zinc oxide and gold, illustrate the behaviour of these, in intimate contact with selenium. The metal-selenium contact yields a poor photocell, the metal-zinc oxide-selenium contact one whose properties are critically dependent on the thickness of the intermediate oxide layer, and the *N* type cadmium oxide-selenium contact one for which the efficiency is high, and the thickness of cadmium oxide not critical. It is suggested therefore that in the practical photocell, the essential mechanism is a contact between two suitable semi-conductors of dissimilar types, any extra metal film when present serving simply to raise the lateral conductivity of the intermediate semi-conducting film when this is not high enough to eliminate undesirable effects of a high internal resistance in the finished cell.

1. INTRODUCTION

In the selenium rectifier (or barrier-layer) photocell, incident radiant energy is transformed into electrical energy in the region of contact between the selenium layer and a thin translucent top-film of some other material. The conversion process clearly depends, in efficiency and in other respects, upon the structure and composition of this contact region. These are difficult to ascertain with certainty, so that as with photo-emissive cells, the technique of making good cells has outstripped knowledge of their precise constitution. This leeway in physical knowledge is all the greater because of the reticence concerning commercial manufacturing processes. The aim of the present work was to dispel some of this uncertainty about the cell's structure, so that structure and theory might be related with less empiricism than hitherto, and the cause of some practical defects such as the Hamaker-Beezhold effect (Hamaker & Beezhold 1934; Bergman & Pelz 1937; Preston 1948) brought to light. The method chosen was to try to develop, in the laboratory, a technique which would yield a cell having characteristics comparable with those of the commercial article, and which would also give the greatest possible information and certainty about the structure. Of physical characteristics the most significant are probably the spectral response and maximum quantum efficiency of the cell, and the maximum open-circuit voltage which it will develop at high illuminations (in clear sunlight, for example). These last may be as high as 0.5 and 0.5 V, respectively, for commercial cells. They all are clearly dependent upon the optical absorption

in the selenium (and probably also in the film) in the operative contact region, and upon the potential barrier developed there. These characteristics are therefore given prominence in the present work.

In practice, of course, the thin film should not absorb much of the light passing through it on the way to the contact region, for this absorption results directly in loss of efficiency, as does also loss of light by reflexion at the cell surface. Further, the electrical resistance of the selenium layer, and the transverse resistance of the film, should be as low as possible, to minimize ohmic losses of energy between the generating contact and the cell terminals. These questions receive attention, but they are not the main concern of the present paper.

2. AVAILABLE INFORMATION AND CHOICE OF MATERIALS

Information available on the selenium cell up to 1934 was summarized in a paper by the author (Preston 1936). Pure selenium, in the annealed crystalline state, upon a base-plate of finely etched or sandblasted steel, was known to give satisfactorily consistent results as a first layer, since variability could be traced with some certainty to the films, particularly to the methods of deposition. Since about 1936, Barnard (1935, 1939) and various Russian workers have observed changes in spectral response due to various additives to the selenium. Generally, these changes were of the kind one would expect to result from the effects of the additives upon the spectral absorption of the layer material. Borelius (1949) has made a study of the physical states and properties of selenium. This sheds some light on the changes which occur in the preparation of the selenium in a photocell (as in the process described later in this paper), but it does not hint at any fundamentally new process.

Thus, failing good reasons to the contrary, selenium in the pure state was chosen as the basic cell material for the present work.

On the thin film much work has been done since about 1936, notably by Görlich (1935, 1939), Görlich & Lang (1938), Eckart & Schmidt (1941), Eckart & Gudden (1941), and Sandström (1946), but with singularly little success in elucidating its precise constitution as distinct from the materials used to produce it. Sandström remarks that 'as regards selenium photo-elements, the chemical composition of the blocking layer has never been finally ascertained'. However, the recent work provides certain pointers. One is that, no matter what the composition of the upper layers of the film, an oxygen or oxide layer in immediate contact with the selenium is favourable to an appreciable voltage output. Sandström assumes that the layer consists most probably of selenium dioxide. This conclusion is consistent with, and largely based upon, the well-known observation that a precious-metal film generally gives better results when deposited by sputtering than when evaporated *in vacuo*. Sputtering in air of course favours the occlusion of an oxygen or oxide layer, without however affording in this case much control over its nature or thickness. Another pointer, from Görlich's work on the influence of the film material on the spectral response, is that the presence of cadmium in the film can be of significance for a good open-circuit voltage and efficiency.

In the present work, therefore, it was recognized that experiments would probably have to cover a fair range of film materials, the two pointers just mentioned being a general guide to their selection. The various techniques used in the work are described in considerable detail, it being considered specially important, in this field, that there should be no impediment to their use, or to confirmation of the results by other investigators. The first process to be described is the preparation of the selenium 'blanks'.

3. PREPARATION OF SELENIUM-COATED BLANKS

The raw materials used for the selenium-coated blanks were 45 mm. diameter disks of very flat, finely sandblasted sheet steel, and selenium twice distilled *in vacuo* from an iron boiler on to a chromium-plated receiver. The heating processes were carried out on a steel hotplate of size $60 \times 10 \times 1\frac{1}{2}$ cm. heated at one end by a gas-ring, so as to provide a suitable range of temperatures. First, enough selenium was melted on to a disk to form a shallow pool, and this was spread with an iron spatula. Then, on to this pool a previously heated and still hot disk of Hysil glass was lowered, in such a way that the selenium 'wetted' the glass without occlusion of bubbles. The resulting sandwich was then quickly removed from the hotplate with forceps and placed *glass downwards* on a flat, cold metal block, and pressed uniformly, so as to expel all but a thin uniform layer of selenium before the selenium solidified in the amorphous state. When quite cold the sandwich was replaced glass downward on the hotplate in a region just below the melting point of selenium. Crystallization then took place in a few minutes, the exuding selenium changing from glassy black to mat grey. On completion of this stage, the sandwich was chilled exactly as before, the glass parted from the crystalline selenium, and the coated steel disk replaced, selenium uppermost, to anneal on the hotplate at a temperature just below the melting point (c. 200°C) for about half an hour.

Certain precautions are necessary in this process. First, neither the selenium nor the glass disk should be too hot in the first stage, but the glass must be hot enough not to produce any crystallization of the selenium. Such crystallization can be seen through the glass, after chilling, as a green patchiness in the otherwise clear black appearance of the adhering selenium. Formation of a reddish film (of condensed selenium crystals) on the glass while it is being lowered on to the molten selenium seems to be advantageous, however, in promoting *subsequent* rapid crystallization. Secondly, the chilling on the metal block should be as severe as possible, and the Hysil disk should not therefore be more than 2 mm. thick. This assists rapid crystallization, with a fine grain size, on reheating.

The selenium thus treated is almost certainly in the granular or fibrous form described by Borelius (1949), not the spherulitic. This form is preferable as having the higher electrical conductivity. Moreover, in so far as variations in the finished photocells could be traced back to the processing of the selenium, it was concluded that the most rapid crystallization gave the best results. The annealing process seems simply to ensure completion of crystal formation in any small residue of amorphous selenium. As a check on the resistance of the finished blanks some

measurements were made. The selenium surface was spray-coated with Wood's metal as a counter electrode. The resulting sandwich acted as a rectifier. However, as the voltage on a rectifier is raised, in the 'through' direction, the differential of voltage with respect to current falls, and finally tends to a constant value, which was found in this case to be about $80\ \Omega$. A similar test on a commercial cell stripped of lacquer gave a value of $50\ \Omega$. These figures represent the resistance of the selenium layer. The comparison seemed satisfactory.

4. PRELIMINARY WORK ON THE THIN FILM

As a first step in the study of the films it was decided to deposit a number of substances on to the blanks by evaporation in a high vacuum, under the cleanest possible conditions. Almost the whole of this preliminary work was done at the H. H. Wills Physical Laboratory of Bristol University during a sojourn of three or four weeks in January 1947, in collaboration with Dr J. W. Mitchell of that Laboratory.

The selenium blanks were provided with suitably disposed sprayed metal contacts enabling the photosensitivity, the transverse film resistance, and the total resistance of the cell to be measured without admission of air to the vacuum chamber. The substances evaporated on to the blanks included gold, cadmium, thallium and aluminium, and a range of compounds of cadmium and thallium with oxygen, sulphur, selenium and tellurium. In most cases, the substance was evaporated from a molybdenum boat.

The observations may be summarized as follows. With gold films, good film conductivity was obtained, but little or no photosensitivity. Also, the cells showed only very slight rectifying properties, or electrical asymmetry as regards back-to-front resistance. With aluminium films, similar results were obtained immediately after deposition of the film, but on admission of air to the chamber both the photosensitivity and electrical asymmetry rose. The former passed through a maximum, and fell again to a very low value on continued exposure to air. The latter continued to increase, though the resistance of the cell in *both* directions eventually became quite large. The transverse conductivity of the film itself remained good, however. With thallium, appreciable photosensitivity and asymmetry were observed very soon after deposition of the film, but the cells deteriorated rapidly—to some extent even before admission of air. Thallic oxide gave quite a high sensitivity *during* the process of evaporation, but this dropped to zero immediately the evaporator was shut off.

Cadmium, cadmium oxide and cadmium oxysulphide all gave a significant photosensitivity, with a notable open-circuit voltage up to about 0.3 V , and appreciable asymmetry of electrical resistance. These materials, however, all underwent some oxidation or dissociation on evaporation, so that the resulting films could probably be assumed to contain a mixture of cadmium and its oxide, with or without sulphur.

The remaining compounds of cadmium and thallium all gave films with a very high electrical resistance, which precluded observation of any photoelectric effect.

To some of these a thin final layer of gold was added, and some sensitivity was obtained in certain cases. The behaviour of these, though not examined further in much detail, suggested that with a double film consisting of an *insulator* surmounted by gold, the thickness of the insulator for optimum sensitivity was very critical. It was also found that the current sensitivity of some of the successful cells could be still further improved by the addition of a 'flash' of gold—a result in accord with references in the literature by Görlich and others. It appeared in this case, however, that the improvement was simply the result of lowering the resistance of the film, for no increase in open-circuit voltage was observed, and so, presumably, the addition of the gold had no influence on the operative contact or 'barrier-layer' in the cell.

The conclusions of this preliminary work could be explained and summarized as follows:

(1) A film of gold (non-reactive metal), deposited directly on selenium, gives only a very low potential barrier, so that rectifying properties are very poor, and the potential difference developed on illumination is too small to drive any appreciable fraction of the photoelectric yield round an external circuit and through the internal resistance of the cell itself.

(2) A film of aluminium (reactive metal), deposited directly on selenium, behaves at first like gold. Later, in contact with moist air, the film combines with the selenium to form an intermediate thin layer of steadily increasing thickness. An appreciable potential barrier, and with it a photocurrent in the external circuit, then develop. The photocurrent falls off again when the intermediate layer exceeds some critical thickness, the potential barrier then presumably becoming so high that the photoelectrons are not sufficiently energetic to penetrate or surmount it.

The intermediate layer here probably consists of (moist?) aluminium selenide, for passage of a current through the cell from an external source then produces a pronounced odour of hydrogen selenide.

(3) Single mixed films consisting probably of cadmium oxide with excess of cadmium give contact conditions favourable to high photosensitivity while having in themselves a more or less adequate electrical conductivity. Owing, however, to the variability of composition obviously resulting from deposition in a high vacuum, no consistent control of the product was obtainable. These results hinted strongly that the study of single films of *semi-conductors*, deposited under controllable conditions, might be particularly profitable.

(4) The remaining experimental results suggested further experiments with double, metal-insulator, films, with varying thicknesses of the insulating component, for the sake of comparison.

Thus the simple metal-selenium contact was practically eliminated from the field, and attention drawn particularly to films of a non-metallic or pseudo-metallic type. The general observations strongly suggested that the optimum system might be selenium in contact with another semi-conductor (perhaps most probably of the *N* type), a second top-film of a metal being added only if the semi-conductor film had too low a conductivity for *practical* requirements. The importance of cadmium oxide as a material for, or at least a component of, the first film, was clearly established.

As the next step in the investigation, therefore, it was decided to experiment further with cadmium and cadmium oxide. Sputtering of cadmium, with subsequent oxidation of the film as necessary (Bädeker 1907), was chosen as the method of deposition—not because sputtering had been found by earlier workers to give better results than evaporation, where simple metal films were concerned, but because it promised to give much better control and to avoid the troubles due to dissociation, etc., which had arisen in the evaporation experiments.

5. THE SPUTTERING PROCESS AND FILMS OBTAINED

A sputtering technique rather different from many previously described, by Ditchburn (1933) for example, was chosen. As the results obtained were unexpected, and had a considerable influence on the course of the work, certain essential details must be noted. The design of the sputtering chamber is shown in figure 1. A substantial copper block was provided to absorb heat generated at the cathode, while all surfaces of the cathode, except that from which deposition was to take place, were screened from the discharge by glass or mica. Provision was made for admitting a continuous stream of argon at a rate sufficient to maintain the correct low pressure while the vacuum pump (a two-stage Speedivac) was running. The surface to receive the sputtered film was blocked up to lie 8 to 9 mm. *below the cathode surface*. A metal block was used here, also, for absorption of heat. The flow of argon and the d.c. high-tension supply were adjusted so that with a circular cathode of 87 mm. diameter and a discharge current of 90 to 100 mA, the sharp boundary of the cathode dark space was 5 to 6 mm. *below the cathode surface*. This method of adjusting the gas pressure was found more practical than using a pressure gauge. The current density was thus much greater than is often used, but it seemed definitely advantageous. On the other hand, considerable heat was generated (the discharge voltage being about 800), and it was found advisable not to run the plant continuously for more than 3 min. at a time when depositing films on selenium. Films requiring longer times were built up in stages, with cooling periods between the runs. The procedure was to pump the chamber hard (no discharge at 2000 V) for a minute or so, then turn on the argon, wait a half-minute or so for pressure equalization, and then switch on the sputtering discharge for as long as required.

Preliminary trials with a gold or silver cathode, cleaned with precipitated chalk, gave films (on glass) which, with only 10 sec. sputtering time, or less, had good electrical conductivity. With a cadmium cathode cleaned with medium glass-paper, however, a similar time gave no visible deposit on a glass slip. Increase of the time to several minutes produced a brownish film. With further increase of time the whole gamut of interference colours could be produced, the film behaving as a transparent, non-metallic one. A chance observation then revealed, however, that it had a remarkably high electrical conductivity for a transparent non-metal. In view of the work at Bristol it was of immediate interest to discover whether the film consisted of a cadmium compound, probably oxide deposited from a superficial

film initially on the cathode, or formed by reaction with traces of oxygen present because the chamber was not thoroughly out-gassed. A film deposited on a metal substrate was therefore examined by electron diffraction (reflexion) in the Metallurgy Division, N.P.L., and identified as cadmium oxide. Unexpectedly, therefore, a process had been found for the production of cadmium oxide films in a highly conducting state. Sputtering in air was found to give films having a far lower

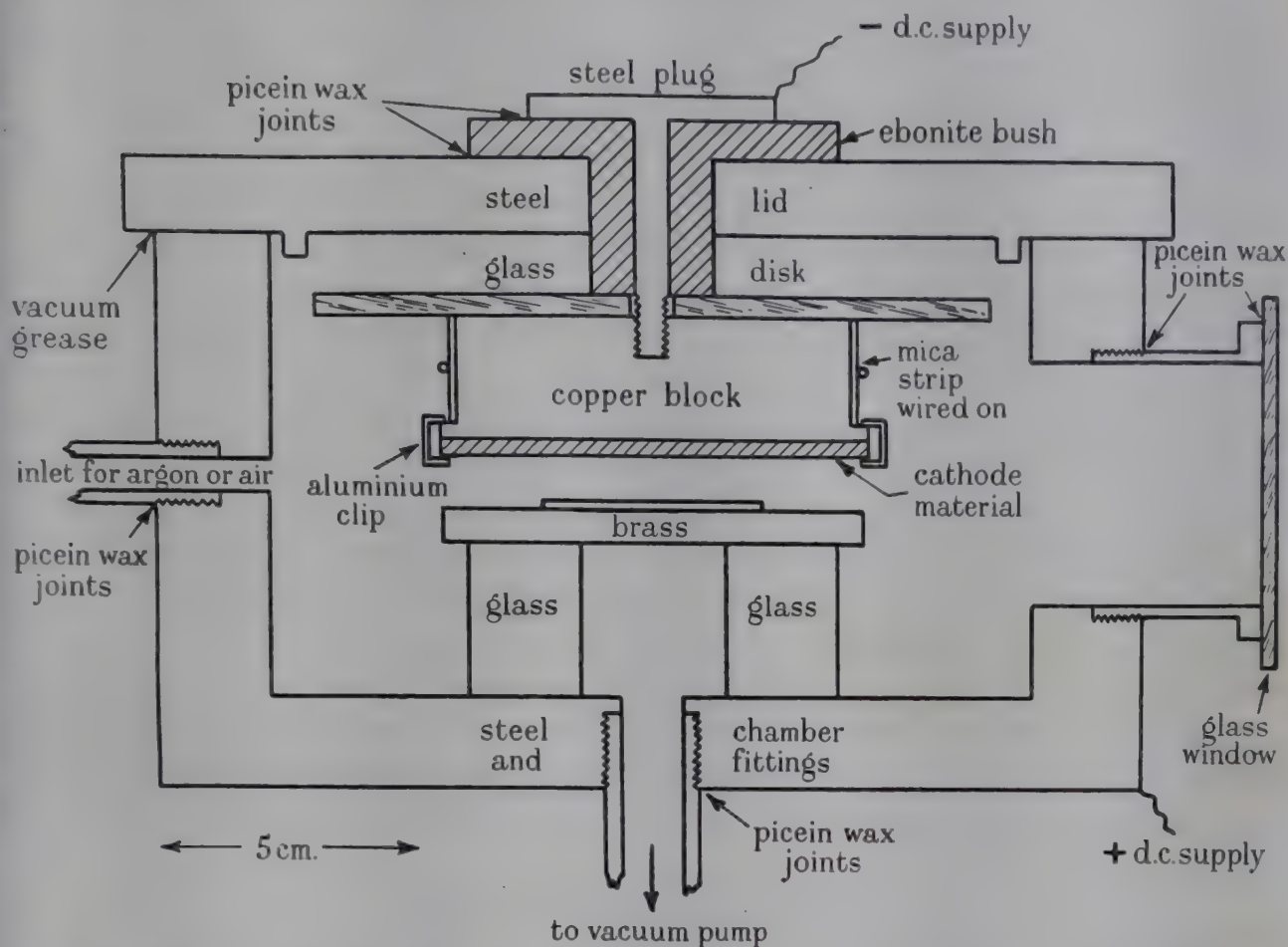


FIGURE 1. Scale drawing of sputtering chamber (vertical in plane of symmetry).
Handles for the lid are not shown.

conductivity. The possibility of producing oxide films by sputtering in oxygen had already been envisaged by Overbeck (1933), who was, however, not concerned with their electrical conductivity. As a matter of interest, a film was sputtered from a zinc cathode and similarly identified as zinc oxide. It did not, however, show a comparably high conductivity. Results with tin were similar.

It was also found that cadmium could be deposited in the metallic state after a prolonged initial pumping of the chamber, or in the presence of a trace of a reducing vapour such as benzene. At this stage it seemed an obvious step to introduce means for controlling the admission of a trace of oxygen to the chamber, along with the argon, and always to start with the plant thoroughly de-gassed. However, a few blank runs quickly brought the plant to a condition in which a high film conductivity was obtained without obvious deposition of metallic cadmium, so this step was not taken.

6. PROPERTIES OF SPUTTERED CADMIUM OXIDE FILMS

Nature of the conductivity

The sputtered cadmium oxide films were then studied in detail. That the oxide was a semi-conductor seemed evident. The readiness with which cadmium oxide dissociated in the evaporation experiments suggested that it should be easy to produce either lattice defects or departure from stoichiometric composition. Earlier work, also, such as that of von Baumbach & Wagner (1933), had shown how variable the conductivity of the bulk material could be, according to conditions of preparation and measurement.

The electron diffraction test, though not a very sensitive one, had revealed a preferred crystal orientation but not any impurity or abnormal composition. A chemical micro-analysis was made of several specimens of film, weighing up to 13 mg. and deposited on clean platinum foil. It was estimated that a departure from stoichiometric composition equivalent to a few per cent of the cadmium content could have been detected, but no certain departure was found by this method. Investigation of the sign of the Hall coefficient, however, gave conclusive evidence that the films consisted of an *N*-type semi-conductor. No quantitative measurement of the coefficient was made, but it was presumably large, for the Hall effect was easily observed using a film-strip 1 cm. wide, carrying a current of 25 mA, in the magnetic field of an ordinary horseshoe instrument magnet, the potential difference developed between the edges of the film being about 100 μ V and easily detected with an ordinary galvanometer. The estimated current density in the film was about 500 A/cm.². This confirmation of the negative sign of the Hall coefficient for the film material is consistent with von Baumbach & Wagner's work (1933) on the bulk oxide, and with later work, such as that of Andrews (1947), on its thermoelectric power.

Thickness and specific resistance

Figure 2 shows spectral reflexion and transmission curves for a fairly thick film, on a glass slip, as taken on the Hardy recording spectrophotometer with light incident near the normal. The periodic variations are of course due to optical interference in the film. The optical path differences in the film, for the reflexion maxima and minima, were identified from a study of a series of thinner films, and are marked in the figure. This study also showed that the refractive index of the film was greater than that of the glass base, so that a phase reversal occurs on reflexion at the air-oxide, but not at the oxide-glass, interface. The approximate thickness was determined by dissolving off part of the film in acid, silvering the step so produced, and studying it in an interferometer. Then using the equation $2\mu t = \frac{1}{2}n\lambda$, where μ = refractive index, t = thickness, and n has the values marked in figure 2, five values of μ as given in table 1 were obtained. The spectral transmission curve shows a strong absorption towards the blue, consistent with the increasing dispersion indicated in the results of table 1.

The thicknesses of films subsequently referred to in this paper were calculated from pairs of values of n and λ taken from measurements of their spectral reflexion, and values of μ taken from a graph of the results of table 1. The assumption was

made, therefore, that all the films had the same refractive indices as given for the first in table 1. This was probably not strictly the case, but it enabled approximate values of thickness to be obtained without destructive operations on the films.

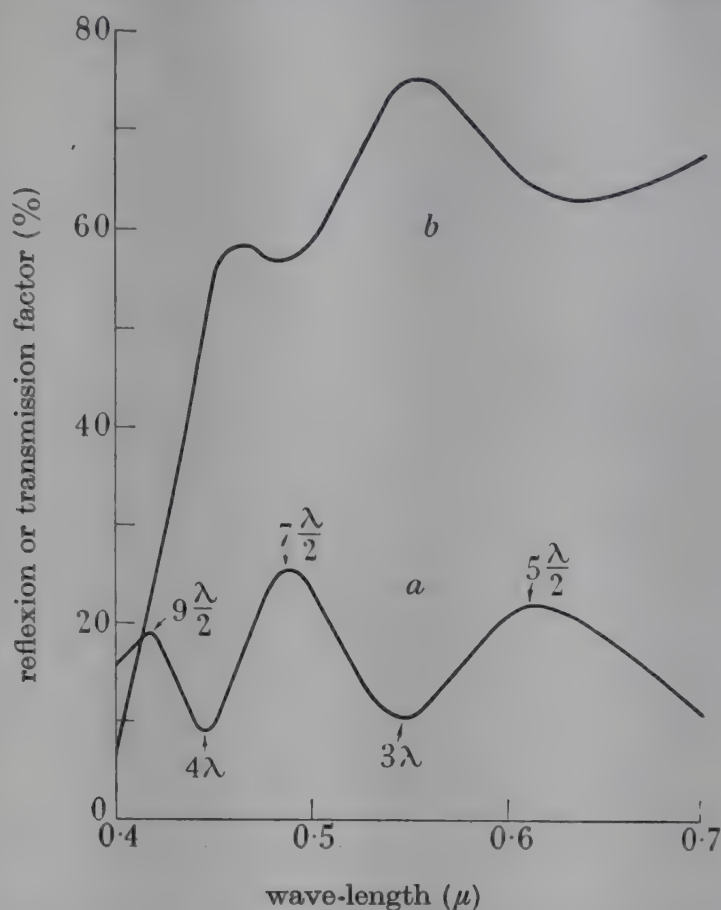


FIGURE 2. Spectral reflexion (*a*) and transmission (*b*) curves of cadmium oxide film on glass for nearly normal incidence (glass included).

TABLE 1. MEASURED THICKNESS $t = 3820 \text{ \AA}$

$2\mu t$	$\lambda (\text{\AA})$	μ
$9\frac{\lambda}{2}$	4170	2.46
$8\frac{\lambda}{2}$	4460	2.34
$7\frac{\lambda}{2}$	4880	2.24
$6\frac{\lambda}{2}$	5440	2.14
$5\frac{\lambda}{2}$	6150	2.01

The specific resistance of two films, processed for high conductivity, was measured with the following results:

	film A	film B
sputtering time (min.)	14	6
thickness ($\text{\AA}$)	2670	1160
length $\times$ breadth (cm.)	3×4	3×4
measured resistance (Ω)	50	80
specific resistance ($\Omega \text{ cm.}$)	1.78×10^{-3}	1.24×10^{-3}

It is remarkable that the resistivity is lower than that of graphite. Also, the values obtained were quite near to those found by Bädcker (1907) for films produced by oxidation of metallic cadmium *after* deposition by sputtering. Some typical sputtered zinc oxide films were also tested and found to have a specific resistance of the order of $10\ \Omega\ \text{cm}$.

Behaviour of the films on photocells

When conducting cadmium oxide films were deposited on the prepared selenium blanks, good photosensitivity was at once obtained, and open-circuit voltages up to about 0.5 V under high illumination were observed. The spectral response, also, was found to be generally similar to that of the average commercial cell, and pronounced rectifying properties were observed. A sputtering time of 3 min. (film thickness 600 to 700 Å) gave a good cell. The reflected colour, owing to interference in the film, was bluish-magenta (first-order interference). A film of about three times this thickness, also of magenta colour (second order), gave rather better results at the higher illuminations, owing no doubt to its lower resistance. It was noted that the reflected colour for a given thickness was complementary to that for a film on glass, selenium having a higher refractive index than the film, and glass a lower. The apparent intensity of the colours suggested that the interference was nearly complete, so that the resulting variation in reflexion through the spectrum might appreciably influence the spectral response of the cell.

Addition of a 'flash' of gold to these cells did not change their properties significantly, except for the small decrease in sensitivity due to optical absorption in the gold. With thinner oxide films, having inadequate conductivity to give a practically useful cell, the sensitivity could be restored by the superposition of a thin gold film. Again, the only difference between such a cell, and one with a thicker film of the oxide only, was attributable to the slightly different optical absorption of the double film. In particular, the open-circuit voltage was practically unchanged. There was, however, a minimum permissible thickness for the oxide in such a double film, below which the high open-circuit voltage was not maintained, but diminished with the thickness of the oxide. This thickness was about 150 Å. This suggests that the space-charge region at the contact between the selenium and the oxide extends into the oxide to a distance of this order, namely, $10^{-6}\ \text{cm}$.—a result in fair agreement with estimates made on theoretical grounds by Mott and others (see Lovell 1947 for references) for typical semi-conductors. Estimates based on measurement of the capacitance of cells have usually led to a value one or two orders of magnitude greater. One might perhaps expect this result, however, when a top metallic film is present, for then the method might give the total thickness of the intermediate layer, acting as a dielectric, rather than that of the space-charge region only.

7. CHARACTERISTICS OF THE PHOTOCELLS

Three cells were chosen to typify the various characteristics. They had films of cadmium oxide only, with upper contact rings of Wood's metal, and were not lacquered. They were:

Cell no. 16. Sputtered for $3\frac{1}{4}$ min. Film thickness about 770 Å. Colour blue. Normal product.

Cell no. 14. Sputtered for 12 min. Film thickness about 2080 Å. Colour bluish-magenta.

Cell no. 15. Sputtered for 15 min. Film thickness about 2620 Å. Colour yellow-green, spectral reflexion roughly complementary to that of no. 14, by intention.

The sensitive area was 11.3 cm^2 in all three cells. Table 2 summarizes sensitivity measurements with white light (tungsten, colour temperature 2850° K), and open circuit voltages.

TABLE 2

external circuit resistance (Ω)	illumination (foot-candles)	cell sensitivity ($\mu\text{A}/\text{foot-candle}$)		
		no. 16	no. 14	no. 15
0	10	8.50	8.83	6.90
0	50	7.53	8.26	6.50
0	100	6.65	7.73	6.05
100	10	8.28	8.68	6.70
100	50	6.70	7.68	5.93
100	100	5.58	6.70	5.10
500	10	7.45	7.95	5.98
500	50	4.68	5.30	4.08
500	100	3.25	3.65	2.88
open-circuit voltage in clear sunlight		0.51	0.49	0.45

Sensitivities as high as $9 \mu\text{A}$ per foot-candle, or $740 \mu\text{A}$ per lumen are therefore obtained. Cell no. 15 gives a low result because of its high reflexion factor in the middle of the visible spectrum.

Figure 3 gives spectral reflexion curves for the cells, and for a free selenium surface, all the surfaces being highly specular. Figure 4 gives the relative spectral response curves, for an equal energy spectrum, all reduced to the same maximum. Comparison of the two figures immediately reveals the effect of selective reflexion on the spectral response. It can be shown to be very nearly what would be expected quantitatively, allowance being made for differences in absorption in the different films.

It is interesting to compute the reflexion factors of the cells, at points of minimum reflexion where the calculation is simple, and compare them with the observed values. For this the refractive index of selenium is required, and is calculated by the Fresnel formula from the curve in figure 3. Table 3 shows such a comparison. Here, the computation included the assumption that there was no optical absorption in the body of the film. Near 4000 Å , however, the film absorbs fairly strongly, and this would account for the disagreement between calculated and observed results in the third line of table 3. It is seen that in a suitable thickness the oxide film 'blooms' the selenium surface, having a refractive index close enough to the square root of that of selenium. The absorption of light in the cell is thus markedly increased, and with it the sensitivity. Cell no. 16 provides a good example of this effect, only a few per cent of the incident light being lost by reflexion, in the middle of the visible spectrum.

The absolute sensitivity of the cells at wave-length 5461 Å was determined by direct comparison with a calibrated thermopile. The results, in terms of incident energy, are given in table 4. This table also gives values of quantum efficiency in terms of absorbed energy, calculated on the basis of reflexion factors at 5461 Å taken from figure 3. (In the literature the quantum efficiency for *absorbed* radiation is very seldom referred to, though it is clearly the more fundamental quantity.) These results may be compared with recent work by Billig & Plessner (1949). The figures were obtained with a low cell output and low circuit resistance. Also, the wave-length 5461 Å is very close to that for maximum spectral response on

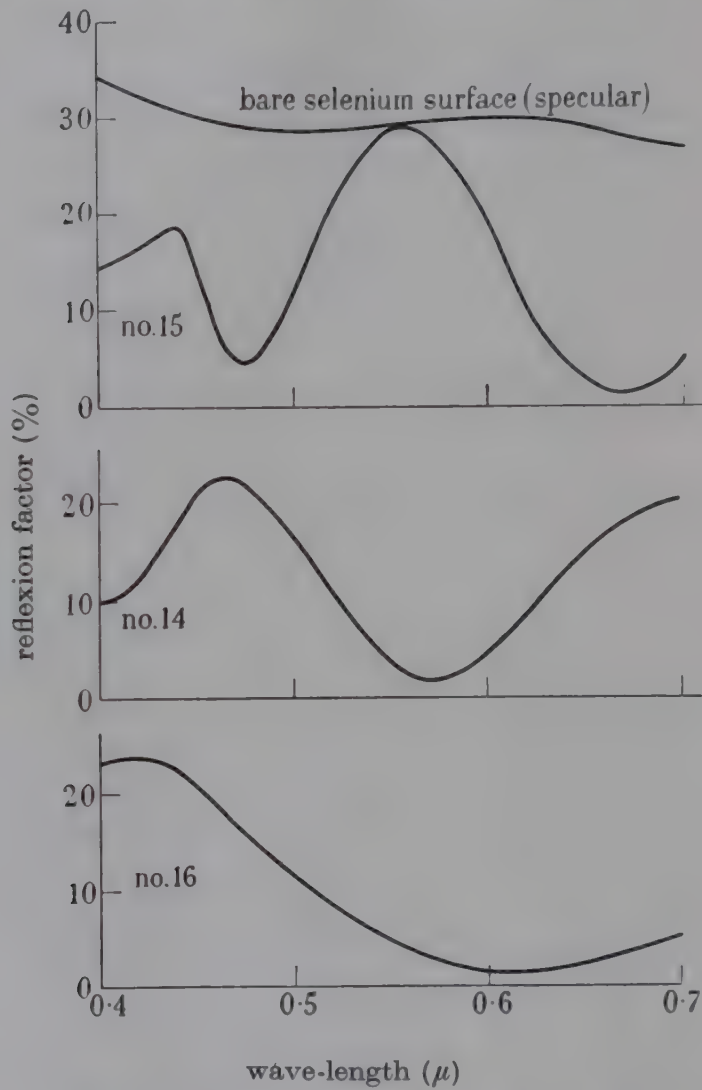


FIGURE 3. Spectral reflexion curves of cells nos. 16, 14 and 15, and of bare selenium surface for nearly normal incidence.

TABLE 3

cell no.	wave-length considered (Å)	reflexion factor of bare selenium	index of selenium	index of film from table 1	reflexion factor of cell	
					calculated	observed
16	6100	0.30	3.42	2.03	0.010	0.01
14	5700	0.29 ₅	3.38	2.09	0.016	0.015
14	4050	0.33 ₅	3.75	2.52	0.067	0.10
15	4750	0.29	3.33	2.27	0.046	0.04

a quantum basis. The results in the last column can therefore be fairly taken to represent the maximum conversion efficiency, on a quantum basis, of the actual cell mechanism, except for a quite small correction for some loss of light in the body of the film.

TABLE 4

cell no.	sensitivity at 5461 Å (A/W)	quantum efficiency (%)	
		for incident light	for absorbed light
16	0.255	58	61
14	0.296	67	70
15	0.207	47	65

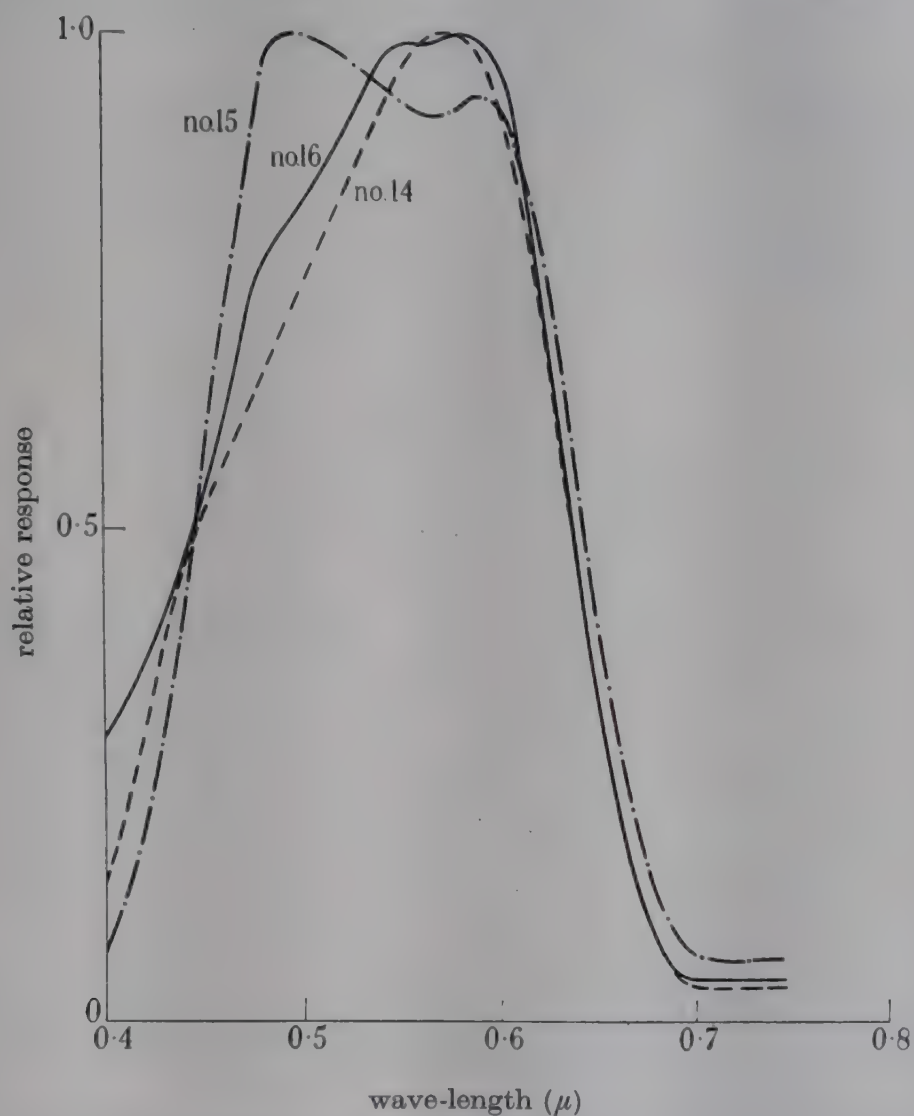


FIGURE 4. Relative spectral response of cells nos. 16, 14 and 15.

Finally, the rectifying properties of the cells were examined. They are shown in the current-voltage curves of figure 5. A number of other cells made in the same way, with various film thicknesses, had characteristics lying very close to, or between, these. The film thickness thus seemed to have very little influence on the rectification characteristic, for all thicknesses greater than was necessary for a good cell (say 200 Å when a top-film of gold was added). Unfortunately, the opportunity was not taken, and has not since occurred, to pursue this study down to smaller thicknesses of cadmium oxide.

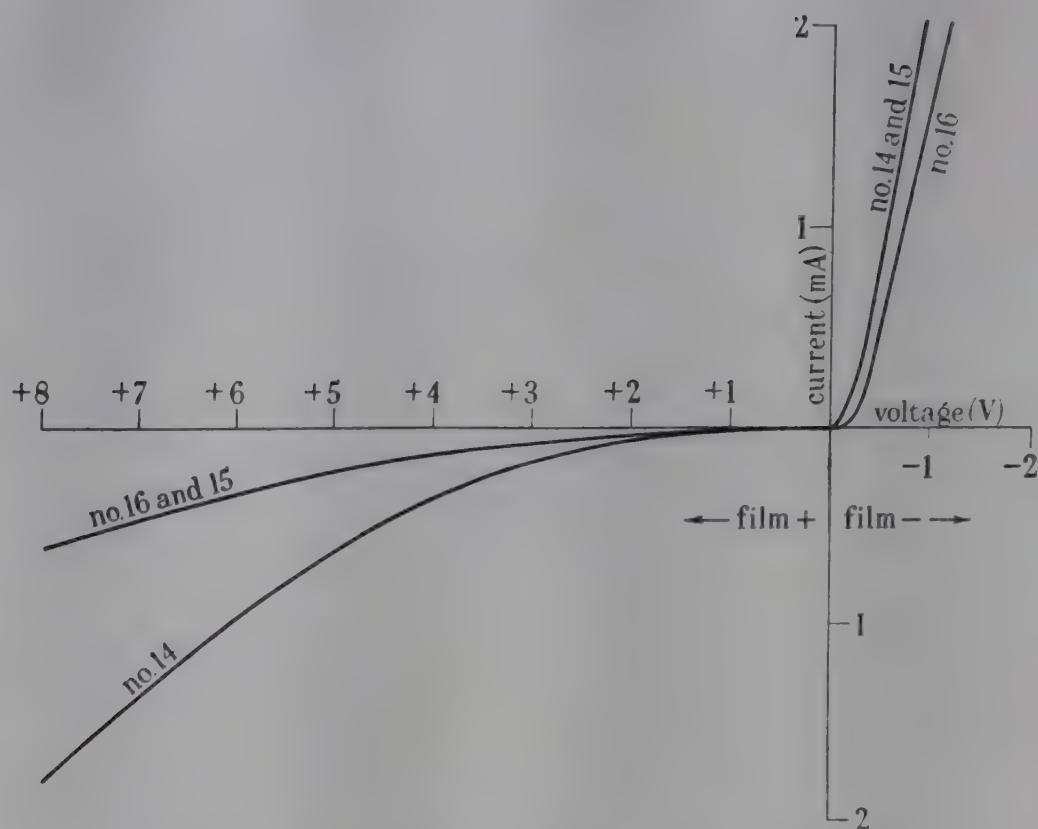


FIGURE 5. Current-voltage curves for typical photocells.

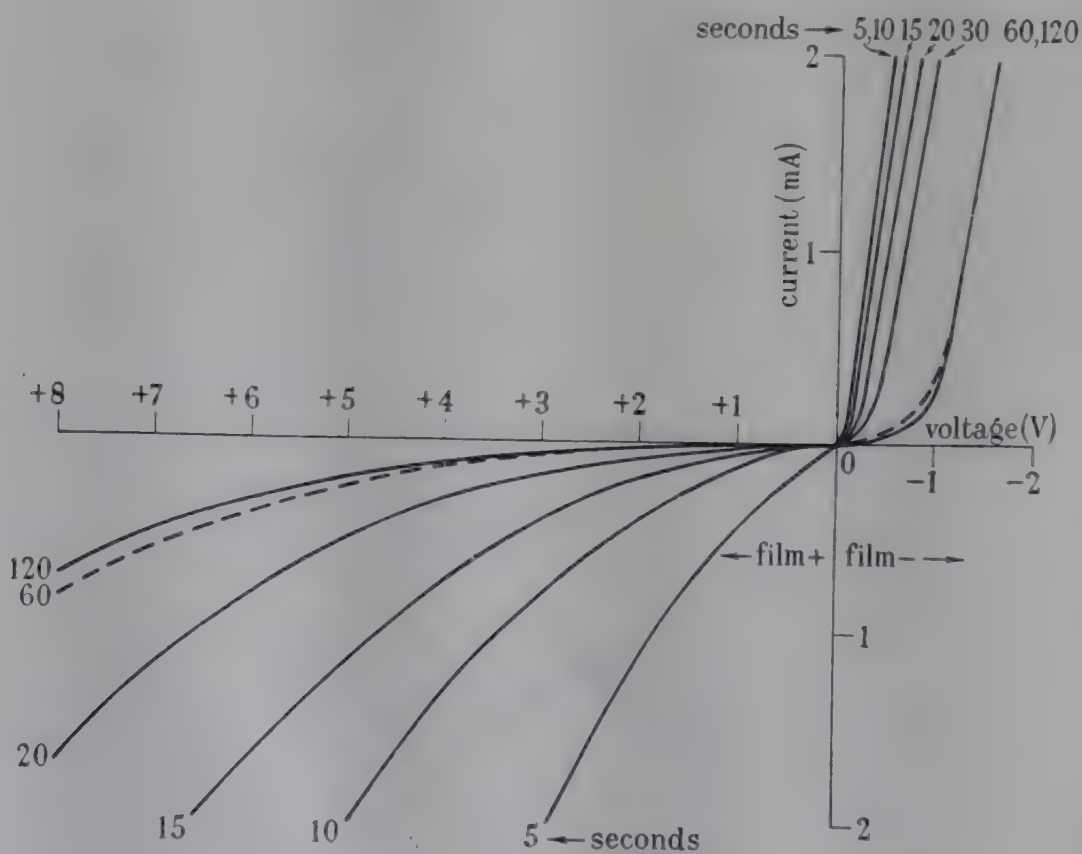


FIGURE 6. Current-voltage curves for disks coated with zinc oxide and gold. (Sputtering time of the oxide shown, in seconds.)

8. EXPERIMENTS WITH ZINC OXIDE FILMS

Experiments were, however, made with very thin zinc oxide films of various thicknesses with a top of gold (7 sec. sputtering time). The observation made earlier that the zinc oxide films had only a low conductivity had suggested that they might be used, in a series of thicknesses, to simulate a selenium-insulator-metal film system. The rectification characteristics of a series of such cells are shown in figure 6. The thicknesses of zinc oxide were not determined directly, but the sputtering times are given instead. On the assumption, only roughly correct, that the zinc oxide builds up at about the same rate as cadmium oxide, the range of thicknesses would be from 350 down to 15 Å. Some photosensitivity was observed, it being greatest with the films of sputtering time 10 and 15 sec., but negligible with films of 1 min. or more sputtering time (thickness greater than 150 Å, roughly). The absolute sensitivity of these cells was at most about one-fifth that of the cadmium oxide cells. The spectral response of the 10 sec. cell is shown in figure 7 as a matter of interest, together with that of a typical cell from a series with tin oxide-gold films which behaved in a similar way. The spectral response was found to vary rather markedly with thickness of the oxide film.

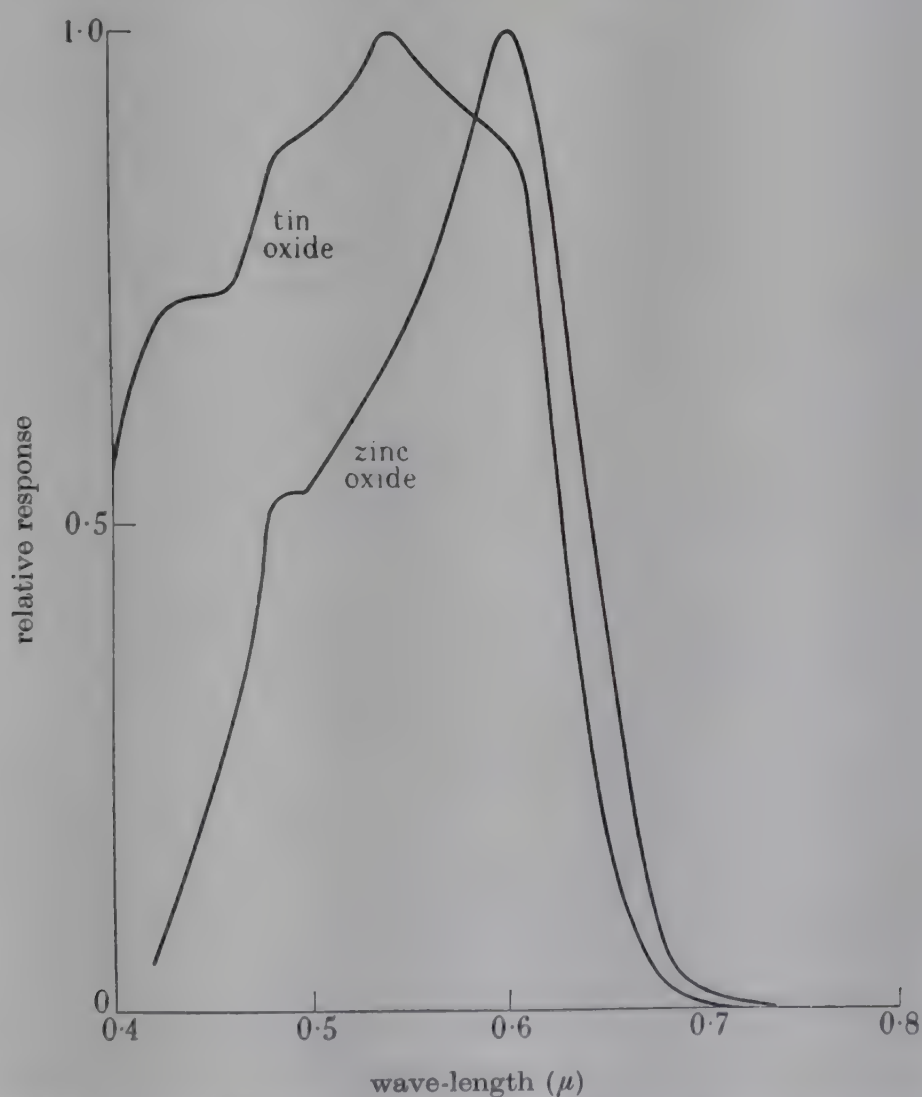


FIGURE 7. Relative spectral response of cells with films of zinc oxide and tin oxide, with gold.

9. DISCUSSION AND CONCLUSIONS

The first aim of the investigation, to establish a controllable technique for making at least one type of selenium rectifier cell which would give a high photo-voltage and quantum yield, was achieved. It is also believed that its structure is defined sufficiently closely. The desired results were obtained with the simple contact between selenium, a *P* type semi-conductor, and cadmium oxide, shown to be an *N* type semi-conductor. It is believed that no significant thickness of any other material was occluded between them. Also the properties of the cells, apart from predictable optical effects, were substantially independent of the oxide thickness from about 10^{-6} cm. up to nearly 3×10^{-5} cm., the greatest thickness used. It is reasonable, then, to take 10^{-6} cm. as roughly the limiting extent of the space-charge region into the oxide, and to regard the contact as an intimate one between two semi-conductors in bulk. The height of the potential barrier would then be characteristic solely of the nature of the two materials. This is the kind of contact, without an inter-layer of a third substance, to which much of Mott's theoretical work has been devoted. The high sensitivity and photo-voltage yielded by the selenium-cadmium oxide contact is no doubt owing to the high optical absorption in the very thin barrier region, together with suitability of the height of the potential barrier naturally developed there. Since superposition of a gold film on the oxide has no fundamentally significant effects, it appears that the additional barrier thereby created at the gold-oxide interface must be very low, as for a selenium-gold contact.

In contrast, the zinc oxide cells probably illustrate a case where the potential barrier at the contact between the selenium and semi-conducting film is not suitable, but too high. In this respect the zinc oxide, though an *N* type semi-conductor, simulates the insulating inter-layer familiar in the work of Schottky. With the thicker oxide films photosensitivity is practically nil. To obtain a good response it is necessary to lower the potential barrier to the proper height by bringing the gold into sufficiently close proximity with the selenium, using an oxide film of controlled thickness less than that of the normal space-charge region in the thicker oxide. The highest sensitivity was in fact obtained with thicknesses around 30 Å. The rectifying properties of the cells, however, change progressively with oxide thickness up to at least 150 Å, which indicates that the space-charge region extends up to this thickness at least, when the oxide is thick enough.

A notable feature of the rectification curves for the zinc oxide cells is that, with the thickest oxide films, the current in the 'through' direction does not rise steeply until the applied voltage approaches 1 V. This suggests an appreciable potential barrier at the oxide-gold interface. From the nature of the materials there in contact, this barrier should be in series-opposition to that at the selenium-oxide contact and would raise the resistance of the cell in the normal 'through' direction in the way observed. Such a reverse barrier would probably be objectionable in a rectifier. The curves show that it is reduced by reducing the oxide thickness, so that the two space-charge regions in the cell overlap, but then the other barrier is also reduced and the 'blocking' resistance of the cell, considered as a rectifier, is lowered. In a cell or a rectifier, therefore, use of an insulator or unsuitable semi-

conductor for the intermediate film would necessitate close control of its thickness to ensure the desired properties.

It seems fairly certain, then, that the essential feature of the constitution of the practical photocell, and probably also the rectifier, is a contact between two suitable semi-conductors of opposite types—in bulk except in so far as admission of light to the contact, or consideration of ohmic losses of energy, may set a practical limitation to the thickness of one or both components.

Secondly, the behaviour of the two- and three-component systems here described may shed light on attempts, like that of Billig & Landsberg (1950), to unify the rectification theories of Mott and Schottky. It should be borne in mind that the thickness of the space-charge region may differ markedly in the different semi-conductors considered. The exceptionally high conductivity of the cadmium oxide suggests (but does not prove) an exceptionally high concentration of impurity centres, which would result in an unusually thin barrier region and steep potential gradient in it.

Thirdly, the work suggests that further practical research of the same kind might be particularly fruitful, attention being given to other materials. Some of those used in the preliminary work, such as thallic oxide, need to be tried again. Others, such as contacts between selenium and *N* type germanium, and between *N* and *P* types of germanium, would seem to merit study. (During the preparation of this paper for press, a sensitive and interesting cell has been made with a sputtered film of lead oxide on selenium, and top film of gold.) Görlich's work on the influence of the film material on spectral response also needs re-examination, and relation to the position of the optical absorption bands of the materials.

Finally, the use of the sputtering technique here described for producing non-metallic films, with close control of thickness, would seem a powerful adjunct to such further practical studies, as well as in the preparation of films for purely optical purposes. In so far as some of the observations in this paper are not novel, they have been repeated because of the added certainty which the technique would seem to give, and to place them alongside the original work done with it.

The author is grateful for collaboration and discussion with staff of the H. H. Wills Physical Laboratory of Bristol University, arranged by kind permission of Professor A. M. Tyndall, and particularly for the generous co-operation of Dr J. W. Mitchell.

Acknowledgements are also due to Messrs Megatron Ltd. of London, who within the limits reasonably to be expected of a commercial firm, placed technical information and small quantities of raw materials at the author's disposal.

In the National Physical Laboratory, Mr E. A. Calnan and Miss I. H. Hadfield of the Metallurgy Division made examinations of the films by electron diffraction and chemical analysis, respectively. The rest of the work was done in the Radiometry Section of the Light Division.

The work described in this paper was carried out under the programme of general research of the National Physical Laboratory, and the paper is published by permission of the Director of the Laboratory.

REFERENCES

- Andrews, J. P. 1947 *Proc. Phys. Soc.* **59**, 990.
 Bädeker, K. 1907 *Ann. Phys., Lpz.*, **22**, 749.
 Barnard, G. P. 1935 *Proc. Phys. Soc.* **47**, 477.
 Barnard, G. P. 1939 *Proc. Phys. Soc.* **51**, 222, and a number of Russian papers.
 Baumbach, H. H. von & Wagner, C. 1933 *Z. phys. Chem.* **22**, 199.
 Bergman, L. & Pelz, R. 1937 *Z. tech. Phys.* **18**, 178.
 Billig, E. & Landsberg, P. T. 1950 *Proc. Phys. Soc. A*, **63**, 101.
 Billig, E. & Plessner, K. W. 1949 *Phil. Mag.* **40**, 568.
 Borelius, G. 1949 *Ark. Fys.* **1**, 305 (includes bibliography).
 Ditchburn, R. W. 1933 *Proc. Roy. Soc. A*, **141**, 169.
 Eckart, F. & Gudden, B. 1941 *Naturwissenschaften*, **29**, 575.
 Eckart, F. & Schmidt, A. 1941 *Z. Phys.* **118**, 199.
 Görlich, P. 1935 *Z. tech. Phys.* **16**, 268.
 Görlich, P. 1939 *Z. Phys.* **112**, 490.
 Görlich, P. & Lang, W. 1938 *Z. phys. Chem. B*, **41**, 23.
 Hamaker, H. C. & Beezhold, W. F. 1934 *Physica*, **1**, 119.
 Lovell, B. 1947 *Electronics*, p. 49 (references). Pilot Press.
 Overbeck, C. J. 1933 *J. Opt. Soc. Amer.* **23**, 109.
 Preston, J. S. 1936 *J. Instn Elect. Engrs*, **79**, 424.
 Preston, J. S. 1948 *Rev. Opt. (théor. instrum.)*, **27**, 520.
 Sandström, A. E. 1946 *Phil. Mag.* **37**, 347.

The Bloch integral equation and electrical conductivity

By P. RHODES, PH.D., *Physics Department, University of Leeds*

(Communicated by E. C. Stoner, F.R.S.—Received 7 February 1950—

Revised 11 April 1950)

An account is given of the solution, for effectively the whole temperature range, of the Bloch (1928) integral equation for the electron momentum distribution in a metal in an electric field. Solutions of this equation, from which the temperature variation of the electrical conductivity of the metal may be immediately calculated, have previously been obtained only in the limiting cases of 'high- and low-temperatures', corresponding to $(T/\theta_D) \gg 1$ and $\ll 1$, where θ_D is the Debye characteristic temperature.

As a preliminary to its solution by numerical methods the integral equation is expressed in a non-dimensional form (§2). Solutions are obtained by deriving a high-temperature approximation which is valid over a much wider temperature range than that previously known, and by means of a method of successive approximations (§3). The temperature variation of conductivity is calculated from these solutions, and it is shown that there are significant differences between the results and those obtained from the semi-empirical formula of Grüneisen (1930) (§4).

A comparison is made between the calculated and observed temperature variation of conductivity for a number of metals. There are deviations in detail, and a brief discussion is given of secondary factors from which they may arise, but in general the agreement is good, and it is concluded that the theoretical treatment covers satisfactorily the main features of the observed variation (§5).

In an appendix it is shown that the approximate relations obtainable by the variational method developed by Kohler (1949) are consistent with the more exact results obtained here.

1. INTRODUCTION

The theoretical treatment of the electrical conductivity of metals, developed by Bloch (1928, 1930), leads to an integral equation for the electron momentum distribution in a metal in an electric field. The temperature variation of conductivity may be estimated directly from the solution of this equation. Up to the present,† however, the Bloch integral equation, or transport equation as it is sometimes called, has been solved only for the special cases of 'high- and low-temperatures', corresponding to $(T/\theta_D) \gg 1$ and $\ll 1$, where θ_D is the Debye characteristic temperature of the metal. In this paper consideration is given to the solution of the transport equation for intermediate temperatures, and so to the estimation of the temperature variation of electrical conductivity over effectively the whole temperature range.

A number of simplifying assumptions are made in the Bloch treatment, and in order to test its applicability to actual metals it is desirable to make as detailed a comparison as possible with experimental results. At high temperatures the theoretical treatment may require modification to take into account the anharmonicity of the lattice vibrations and the effects of thermal expansion; at low temperatures the 'residual' resistance due to impurities and strains may be relatively large, and a large, and uncertain, 'correction' may be required to determine the 'ideal' resistance from that observed. Consequently it is at intermediate temperatures that comparison between theory and experiment would be most significant.

Solution of the transport equation by numerical methods, such as those adopted here, is necessarily carried out for particular values of the parameters involved. By considering the physical significance of these parameters, however, it is possible to select values for them which are appropriate to a wide range of actual metals, and to solve the equation for intermediate temperatures for these values (§3). The associated temperature variation of conductivity may then be calculated (§4) and compared with the experimental results (§5).

Full accounts of the theoretical treatment of electrical conductivity, based largely on Bloch's original papers, are given by Brillouin (1931), Sommerfeld & Bethe (1933), Wilson (1936) and others, and it is unnecessary to go over the long arguments in detail here; but a brief outline of the treatment is essential to introduce the various quantities involved (§2). Reference may also be made to the simplified treatment given by Mott & Jones (1936), which shows in a very direct way the relation between conductivity and temperature; since this treatment is valid only at high temperatures, however, it is not considered in detail here.

2. THE BLOCH TREATMENT OF ELECTRICAL CONDUCTIVITY

In the absence of an applied field the distribution of the electrons in a metal among the available momentum states is determined by the equilibrium, Fermi-Dirac, distribution function

$$f_0(\mathbf{k}) = f_0(\epsilon) = \left\{ \exp \left(\frac{\epsilon - \zeta}{kT} \right) + 1 \right\}^{-1}, \quad (2.1)$$

† After this paper had been completed a paper by Kohler (1949) dealing with the solution of the Bloch integral equation became available; its relation to the present paper is considered in an appendix.

where $\mathbf{k}$ is the electron wave-vector, and ϵ and ζ are respectively the energy and chemical potential per electron. In an applied field the distribution is modified by the field and by 'collisions'. In an electric field parallel to the x -axis the steady-state, non-equilibrium, distribution function, $f(\mathbf{k})$, may be expressed in the form (Wilson 1936, pp. 204, 206)

$$f(\mathbf{k}) = f_0(\mathbf{k}) - k_x \mu(\epsilon) \frac{df_0}{d\epsilon}, \quad (2.2)$$

where k_x is the x -component of $\mathbf{k}$, and $\mu(\epsilon)$ is a function to be determined.

Since the steady-state distribution function is independent of time it satisfies the Boltzmann equation

$$\frac{df}{dt} = \left(\frac{\partial f}{\partial t} \right)_{\text{field}} + \left(\frac{\partial f}{\partial t} \right)_{\text{coll.}} = 0. \quad (2.3)$$

To obtain an explicit equation for f from (2.3), expressions are required for $(\partial f / \partial t)_{\text{field}}$ and $(\partial f / \partial t)_{\text{coll.}}$ in terms of f . A general expression for $(\partial f / \partial t)_{\text{field}}$ may be readily obtained (see, for example, Wilson 1936, p. 63). For $(\partial f / \partial t)_{\text{coll.}}$ the problem is much more complicated, and the approximate expressions derivable are dependent on the particular collision mechanism presumed to be involved. In the treatment of 'collisions' associated with thermal lattice vibrations, the case considered here, two different simplifying assumptions have been adopted (Bloch 1928, cf. Sommerfeld & Bethe 1933; Nordheim 1931, cf. Wilson 1936). Both approximations lead to formally similar results; the essential difference lies in the physical significance of a function, depending on the potential field of the lattice ions and on the electron wave-functions, and usually denoted by C , which occurs in the final expression for $(\partial f / \partial t)_{\text{coll.}}$. As will be seen later, the two treatments lead to the same calculated temperature variation of conductivity, though the calculated absolute values would be different.

In deriving the expression for $(\partial f / \partial t)_{\text{coll.}}$ it is also assumed that the relation between electron energy, ϵ , and wave-vector, $\mathbf{k}$, is of the form

$$\epsilon = \frac{\hbar^2}{8\pi^2 m^*} |\mathbf{k}|^2, \quad (2.4)$$

corresponding to electrons of effective mass m^* distributed in a single band of standard form.

The final expression for $(\partial f / \partial t)_{\text{coll.}}$ is an integral involving the function μ , and by combining this with (2.3) and the expression for $(\partial f / \partial t)_{\text{field}}$ the Bloch integral equation for μ (and hence for f) may be obtained in the forms given by Sommerfeld & Bethe (1933, p. 521, equation (35.19)) and Wilson (1936, p. 215, equation (345); the present case corresponds to the inclusion of only the first term on the left-hand side of that equation). In order to bring out more clearly the dimensional character of the various functions involved the Bloch equation may be conveniently re-expressed in the following non-dimensional form:

$$\chi^\dagger = \int_0^y \{ \Psi(z) - \Psi(-z) \} dz, \quad (2.5)$$

where

$$\left. \begin{aligned} \Psi(z) &= \frac{z^2(e^\xi + 1)}{(1 - e^{-z})(e^{\xi+z} + 1)} \left\{ \left(\chi + \frac{pz}{2r} - \frac{z^2}{y} \right) \phi(\xi + z) - \chi \phi(\xi) \right\}; \\ \phi(\xi \pm z) &= \mu(\epsilon \pm kTz) / \{bMk\theta_D(h/m^*C)^2 y^{\frac{1}{2}} eF\}; \\ \xi &= (\epsilon - \zeta)/kT = x - \eta; \\ y &= \theta_D/T; \quad z = h\nu/kT; \\ r &= \epsilon_0/k\theta_D = \theta_F/\theta_D; \quad p = (2q^2)^{\frac{1}{2}}; \\ \chi &= \frac{p}{r} \frac{\epsilon}{kT} = py \left\{ \frac{\zeta}{\epsilon_0} + \frac{\xi}{ry} \right\}; \\ 1/b &= 3(2^{\frac{1}{2}}\pi)^3. \end{aligned} \right\} \quad (2.6)$$

In these equations the symbols have the following significance: F , applied field; $-e$, electron charge; M , mass of atom; ν , frequency of lattice vibration; θ_D , Debye characteristic temperature; ϵ_0 , Fermi zero energy; q , number of electrons per atom outside full bands. The function C , which is assumed to be independent of electron energy and of temperature, may be taken as equal to that given by Wilson (1936, p. 201) or, apart from a numerical factor of $\frac{2}{3}$, that given by Sommerfeld & Bethe (1933, p. 513). In the derivation of (2.5) use has been made of the relation between ϵ_0 and the number of electrons per unit volume, n (see Stoner 1939, p. 267):

$$\epsilon_0 = \frac{h^2}{8m^*} \left(\frac{3n}{\pi} \right)^{\frac{2}{3}} = \frac{h^2}{8m^*a^2} \left(\frac{3q}{\pi} \right)^{\frac{2}{3}}, \quad (2.7)$$

where a^3 = volume of atomic cell.

The electric current density associated with the steady-state distribution (2.2) is

$$\begin{aligned} J_x &= -(e/4\pi^3) \int v_x f dk_x dk_y dk_z \\ &= \frac{e}{4\pi^3} \int \frac{2\pi}{h} \left(\frac{\partial \epsilon}{\partial k_x} \right) k_x \mu(\epsilon) \frac{df_0}{d\epsilon} dk_x dk_y dk_z, \end{aligned}$$

which with (2.4) may be expressed as

$$J_x = \frac{16\pi^2(2m^*)^{\frac{1}{2}}e}{3h^4} \int_0^\infty \epsilon^{\frac{1}{2}} \mu(\epsilon) \frac{df_0}{d\epsilon} d\epsilon. \quad (2.8)$$

Using this expression and (2.6), (2.7) the conductivity, σ , is given by

$$\left. \begin{aligned} \sigma &= J_x/F = \alpha y \int_{-\eta}^\infty \chi^{\frac{1}{2}} \phi(\xi) \frac{df_0}{d\xi} d\xi, \\ \alpha &= \{Mhk\theta_D e^2\} / \{6a^3(\pi m^*C)^2\}. \end{aligned} \right\} \quad (2.9)$$

where

In order to calculate the absolute conductivity from (2.9) it would be necessary to determine the function C , while to calculate the temperature variation of conductivity it is sufficient to solve the transport equation, (2.5), for $\phi(\xi)$ for a range of values of y .

High-temperature approximation

Bloch (1928) obtained a solution of (2.5) for $y \ll 1$ effectively by taking the first term in a Taylor series expansion of $\phi(\xi \pm z)$ and the first term in the expansion of the

remainder of the integrand in ascending powers of z . The right-hand side of (2.5) then involves the function ϕ only in the form $\phi(\xi)$, and (2.5) is reduced from an integral equation to a linear algebraic equation. Neglecting the small terms $\pm (pz/2r)$ (cf. §3), the solution of this equation is given by

$$\phi(\xi) = -2\chi^1/y^2. \quad (2.10)$$

The corresponding value of the conductivity from (2.9) is

$$\sigma = -\frac{2\alpha}{y^2} \int_{-\infty}^{\infty} \chi^1 \frac{df_0}{d\xi} d\xi. \quad (2.11)$$

The Fermi-Dirac integral in this expression cannot be evaluated analytically, but since in many metals the conduction electrons form a highly degenerate 'electron gas' at ordinary temperatures it is sufficient to use the first term in the 'low-temperature' approximation to the integral. From (2.6), (2.11) and the Fermi-Dirac approximation,

$$\sigma = 2\alpha(py)^3/y^2 = 4\alpha q^2 y. \quad (2.12)$$

This expression indicates that the conductivity at high temperatures varies inversely as the absolute temperature, corresponding to a specific resistance, $\rho_T = 1/\sigma$, proportional to T . Rough calculations show that this relationship is not appreciably affected by using more exact values for the Fermi-Dirac integral in (2.11) provided $(T/\theta_F) < 0.1$.

Low-temperature approximation

Accounts of the method used by Bloch (1930) to obtain an approximate solution of (2.5) for $y \gg 1$ have been given by Sommerfeld & Bethe (1933, p. 526), Wilson (1936, p. 212) and others. The method is also discussed in §3, and it will suffice here to quote the expression obtained for the conductivity.

$$\sigma = \frac{\alpha q^2}{5! \zeta(5)} y^5, \quad (2.13)$$

where $\zeta(5)$ is a Riemann function. Equation (2.13) indicates a variation of specific resistance, ρ_T , with T^5 .

3. SOLUTION OF THE TRANSPORT EQUATION FOR INTERMEDIATE TEMPERATURES

Two methods of obtaining solutions of the transport equation for intermediate temperatures are considered in this section. First, a modified form of high-temperature solution, valid over a much wider temperature range than that previously given, is obtained. Secondly, a numerical method of successive approximations for testing the validity of the new first approximation and for obtaining higher approximations is described. For both of these the expression of the transport equation in the non-dimensional form, (2.5), introduced in §2 was an essential preliminary.

Numerical values of parameters

Equation (2.5) contains parameters, ξ , y , z , which are reduced variables appropriate to any metal, and others, p and r , which may be regarded as specifying a particular metal. The function χ depends upon both kinds of parameter and upon ζ and ϵ_0 .

(2.5) is to be solved numerically, it is desirable to select values of p and r which will give the solution as wide a range of applicability as possible.

For metals in which the electrons satisfy the degeneracy criterion, $(kT/\epsilon_0) \ll 1$, ϵ_0 is closely equal to ϵ_0 , and it is sufficiently accurate to express χ (cf. (2.6)) as

$$\chi = py \left\{ 1 + \frac{\xi}{ry} \right\}. \tag{3.1}$$

In order to deal with electrons in a narrow band, i.e. for which ϵ_0 is comparable with T , it would be necessary to take into account the dependence of ζ on T . This would not involve any insuperable difficulty, since Stoner (1939) has given numerical tables and series expressions which show the dependence of ζ on T for effectively the whole temperature range, but in what follows χ is taken to be of the form (3.1) so as to avoid an inessential complication of the central problem.

The number of electrons per atom, q , outside full bands, which determines p (cf. (2.6)), is taken as equal to unity, so that the results are immediately applicable to the alkali and noble metals (cf. §4).

TABLE 1. DEBYE TEMPERATURES, ESTIMATED FERMI-DIRAC CHARACTERISTIC TEMPERATURES AND RELATED DATA FOR REPRESENTATIVE METALS

A , atomic weight;
 d , density, g.cm.⁻³;
 θ_D , Debye characteristic temperature;
 θ_F , Fermi-Dirac characteristic temperature, ϵ_0/k (see equation (3.2));
 $r = \theta_F/\theta_D$.

	A	d	θ_D	$\theta_F \times 10^{-4}$	$r \times 10^{-2}$
Li	6.94	0.534	(400)	5.46	(1.36)
Na	23.00	0.97	(150)	3.66	(2.44)
K	39.10	0.86	100	2.37	2.37
Rb	85.48	1.52	(70)	2.06	(2.94)
Cs	132.91	1.87	(50)	1.76	(3.52)
Cu	63.57	8.93	315	8.15	2.59
Ag	107.88	10.50	215	6.38	2.97
Au	197.2	19.3	170	6.41	3.77

In table 1 values of r , equal to θ_F/θ_D , and of the data used in estimating θ_F , are given for some representative metals. The Debye temperatures, θ_D , are from Borelius (1935, p. 252) and Seitz (1940, p. 110). Values enclosed in brackets are less certain than the others. The Fermi-Dirac characteristic temperatures, θ_F , have been determined from (2.7), which, for electrons to the number q per atom in a metal of density d , atomic weight A , may be written (cf. Stoner 1939, p. 279)

$$\theta_F = \epsilon_0/k = 3.017 \times 10^5 (m/m^*) (qd/A)^{\frac{1}{3}}, \tag{3.2}$$

using the Birge (1941) values of the fundamental constants. The values are calculated for $q = 1$, $m^* = m$. Values of A and d are from standard tables.

The values of r in table 1 are all of the same order of magnitude, and much greater than unity, the mean value being $2.7_4 \times 10^2$. As will be seen later, the solution of (2.5) for $r \gg 1$ does not depend critically upon r . A value of $r = 3 \times 10^2$ is adopted in the numerical work. This value is sufficiently close to all those listed in table 1 for the results of the calculations to be directly applicable to at least all those metals.

Strictly, $\phi(\xi)$ is to be determined from (2.5) for all values of ξ from $-\eta$ to ∞ (corresponding to values of ϵ from 0 to ∞), but it can be seen from (2.9) that in order to calculate the conductivity accurately it is necessary to obtain $\phi(\xi)$ accurately only for a limited range of values of ξ . From (2.1), (2.6),

$$\frac{df_0}{d\xi} = -\frac{1}{(e^\xi + 1)(e^{-\xi} + 1)}, \quad (3.3)$$

so $|df_0/d\xi|$ decreases exponentially for large values of $|\xi|$. Consequently the only appreciable contribution to the integral in (2.9) comes from the region of small values of $|\xi|$; $|df_0/d\xi|$ decreases to 1 % of its value at $\xi = 0$ when $|\xi| \doteq 6$. In order to obtain an accurate value of σ it is sufficient therefore to obtain a solution of (2.5) of comparable accuracy only for small values of $|\xi|$; for larger values of $|\xi|$ a lower degree of accuracy in the values of $\phi(\xi)$ is adequate.

First approximation

Using the Taylor series expansion for $\phi(\xi \pm z)$ in (2.5), the part of the integrand involving ϕ becomes

$$\left[\left\{ \left(\frac{pz}{2r} - \frac{z^2}{y} \right) \phi(\xi) + z \left(\chi + \frac{pz}{2r} - \frac{z^2}{y} \right) \phi'(\xi) + \dots \right\} \frac{e^z}{e^{\xi+z} + 1} + \left\{ \left(-\frac{pz}{2r} - \frac{z^2}{y} \right) \phi(\xi) - z \left(\chi - \frac{pz}{2r} - \frac{z^2}{y} \right) \phi'(\xi) + \dots \right\} \frac{1}{e^{\xi-z} + 1} \right],$$

where a dash denotes differentiation with respect to ξ . Examination of (2.5) suggests that a modified high-temperature solution, valid over a wider range of y , may be obtained by neglecting $\phi', \phi'', \dots$ and the small terms $\pm (pz/2r)$, but leaving the remainder of the integrand unaltered. The tentative new first approximation to $\phi(\xi)$, say $\phi_1(\xi)$, is then

$$\phi_1(\xi) = -y\chi^{\frac{1}{2}}/I(\xi, y), \quad (3.4)$$

where

$$I(\xi, y) = (e^\xi + 1) \int_0^y \frac{z^4 dz}{e^z - 1} \left\{ \frac{e^z}{e^{\xi+z} + 1} + \frac{1}{e^{\xi-z} + 1} \right\}. \quad (3.5)$$

Values of $I(\xi, y)$, determined by numerical integration and by series expansion, are shown in figure 1 in the reduced form $I(\xi, y)/I(0, y)$ for $y = 1, 2, 3, 5$. The corresponding asymptotic values, determined from

$$I(\infty, y) = \int_0^y \frac{z^4 (e^\xi + 1) dz}{e^z - 1}, \quad (3.5a)$$

are also shown.

For $y < 1$, expansion of the integrand of (3.5) in ascending powers of z and term-by-term integration gives

$$I(\xi, y) = \frac{1}{2}y^4 + \frac{1}{36}\{1 - 12e^\xi(e^\xi + 1)^{-2}\}y^6 + \dots$$

Substituting the first term of this series for I in (3.4), the expression for ϕ_1 becomes the same as the previously known high-temperature approximation, (2.10), which is thus a special case of (3.4) for $y^2 \ll 1$. For larger values of y it is not possible to estimate *a priori* the error introduced in (3.4) by the neglect of $\phi', \phi'', \dots$. The most direct method of determining the range of validity of (3.4) is to substitute $\phi_1(\xi)$ in (2.5) and to test whether or not that equation is satisfied. The details of this method

may best be considered in conjunction with the discussion of the method of successive approximations used to obtain higher approximations in the temperature range for which ϕ_1 is not sufficiently accurate.

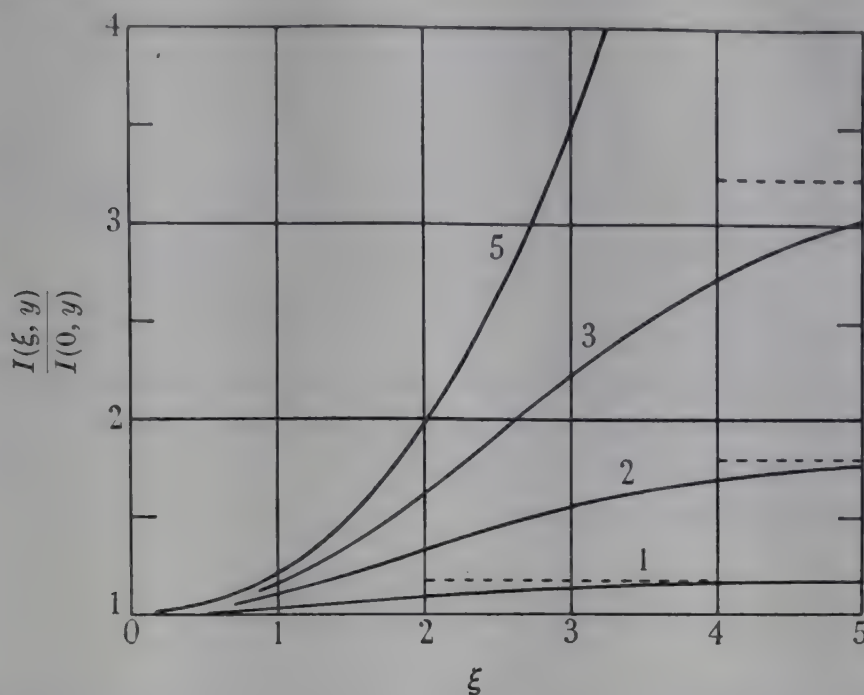


FIGURE 1. The reduced function $I(\xi, y)/I(0, y)$; see equation (3.5). The figures on the curves give the values of y . The values of $I(0, y)$ for $y = 1, 2, 3, 5$ are 0.4490, 5.376, 18.15, 54.15 respectively. The broken lines show the values of $I(\infty, y)/I(0, y)$ for $y = 1, 2, 3$ calculated from (3.5a); the corresponding value for $y = 5$ is 12.07.

Solution by successive approximations

The method of successive approximations adopted may be indicated by rewriting (2.5) in the form

$$\chi K(\xi, y) \phi_{\nu+1}(\xi) = -\chi^{\frac{1}{2}} + \int_0^y \{F(z) - F(-z)\} dz, \quad (3.6)$$

where

$$F(z) = \frac{z^2(e^\xi + 1)}{(1 - e^{-z})(e^{\xi+z} + 1)} \left(\chi + \frac{pz}{2r} - \frac{z^2}{y} \right) \phi_\nu(\xi + z), \quad (3.7)$$

and $K(\xi, y)$ is of the same form as $I(\xi, y)$, (3.5), but with the factor z^4 in the integrand replaced by z^2 . The subscripts to the ϕ 's have been added in (3.6), (3.7) for convenience in discussing the method of successive approximations. The 'solution' of (3.6) is to be understood to mean the solution, $\phi(\xi)$, of the equation obtained by omitting the subscripts.

If a particular function ϕ_ν is a solution of (3.6) then obviously the function $\phi_{\nu+1}$, obtained by inserting ϕ_ν in the right-hand side, will be equal to ϕ_ν . This test for consistency provides a check on any approximate solution which may be obtained. In the present connexion, agreement to within 1 % between $\phi_{\nu+1}$ and ϕ_ν was regarded as sufficient, and this condition need be satisfied only for small values of $|\xi|$, up to 5 say, since the main contribution to the integral in the expression for the conductivity, (2.9), arises from such values of ξ .

In order to obtain higher approximations in the range of y values for which ϕ_1 does not satisfy the consistency condition the following method has been used. Values of

ϕ_2 are determined from (3.6) in the process of checking ϕ_1 , and values of ϕ_3 are determined in the same way from ϕ_2 . If ϕ_3 agrees with ϕ_2 to within the required accuracy for the relevant values of ξ , then ϕ_2 is a satisfactory solution; if not, the process is repeated until a function, ϕ_ν , is obtained such that $\phi_{\nu+1}$ agrees with ϕ_ν to within the required limits. In practice this procedure involves extensive computational work, since, for each value of y , numerical integration of the integral on the right-hand side of (3.6) gives a value of $\phi_{\nu+1}$ for only one value of ξ , and computation of $\phi_{\nu+2}$ from $\phi_{\nu+1}$ requires $\phi_{\nu+1}$ for a whole range of values of ξ .

Examination of (3.6) indicates that the solution is such that $\phi(-\xi)$ is approximately equal to $\phi(\xi)$. In obtaining successive approximations it was assumed that ϕ_ν is an even function at all stages. The final approximation was then tested for consistency in (3.6) for both negative and positive values of ξ , and in all cases it was found that the final value of $\phi_{\nu+1}(-\xi)$ agreed to within 1 % with $\phi_{\nu+1}(\xi)$. This considerably reduced the necessary computation.

From (3.4), (3.1), ϕ_1 may be written as

$$\phi_1(\xi) = - \frac{(p^3 y^5)^{\frac{1}{2}} \{1 + (\xi/ry)\}^{\frac{1}{2}}}{I(\xi, y)}. \quad (3.8)$$

For small values of ξ , and values of y of order unity, $(\xi/ry) \ll 1$. Consequently from (3.8) $\phi'_1, \phi''_1, \dots$ are determined largely by the corresponding differential coefficients of I . Further, the method of derivation indicates that ϕ_1 is a good approximation to ϕ if $\phi', \phi'', \dots$ are small compared with ϕ . The function ϕ_1 may therefore be expected to be a good approximation to the solution of (3.6) if the differential coefficients of I with respect to ξ are small compared with I . The rate of change of I with ξ increases with y (cf. figure 1), so if ϕ_1 is a satisfactory approximation for a particular value of y it may be concluded that it is also satisfactory for smaller values of y .

For $y = 1$ it is found that ϕ_1 is a sufficiently accurate approximation, since ϕ_2 agrees with ϕ_1 to within 1 % for $|\xi| \leq 5$. For $y = 2$ higher approximations are required. In order to determine the largest value of y for which ϕ_1 is a sufficiently accurate solution it would be necessary to test the validity of this solution for a number of values of y in the range 1 to 2. This is unnecessary, however, since in the calculation of the conductivity associated with ϕ it is found that for $y \leq 2$ it is sufficient to consider only ϕ_1 (see §4).

Detailed application of the method of successive approximations has been restricted to $y = 2, 3, 5$ (cf. §§4, 5). The shape of the curves representing the successive approximations are similar for all these values of y , but more approximations are required as y increases. In illustration the successive approximations, ϕ_ν , to the solution of (3.6) for $y = 5$ are shown in figure 2. The curve showing ϕ_4 , which lies between those for ϕ_3, ϕ_5 , has been omitted for clarity. The values of ϕ_6 , indicated by circles, agree to within 1 % with the corresponding values of ϕ_5 , which is taken as the final approximation.

For $y = 2$ and 3 it is the second and third approximations, respectively, which constitute sufficiently accurate approximations to the solution of (3.6). In what follows it will be convenient to denote this final approximation in each case by ϕ .

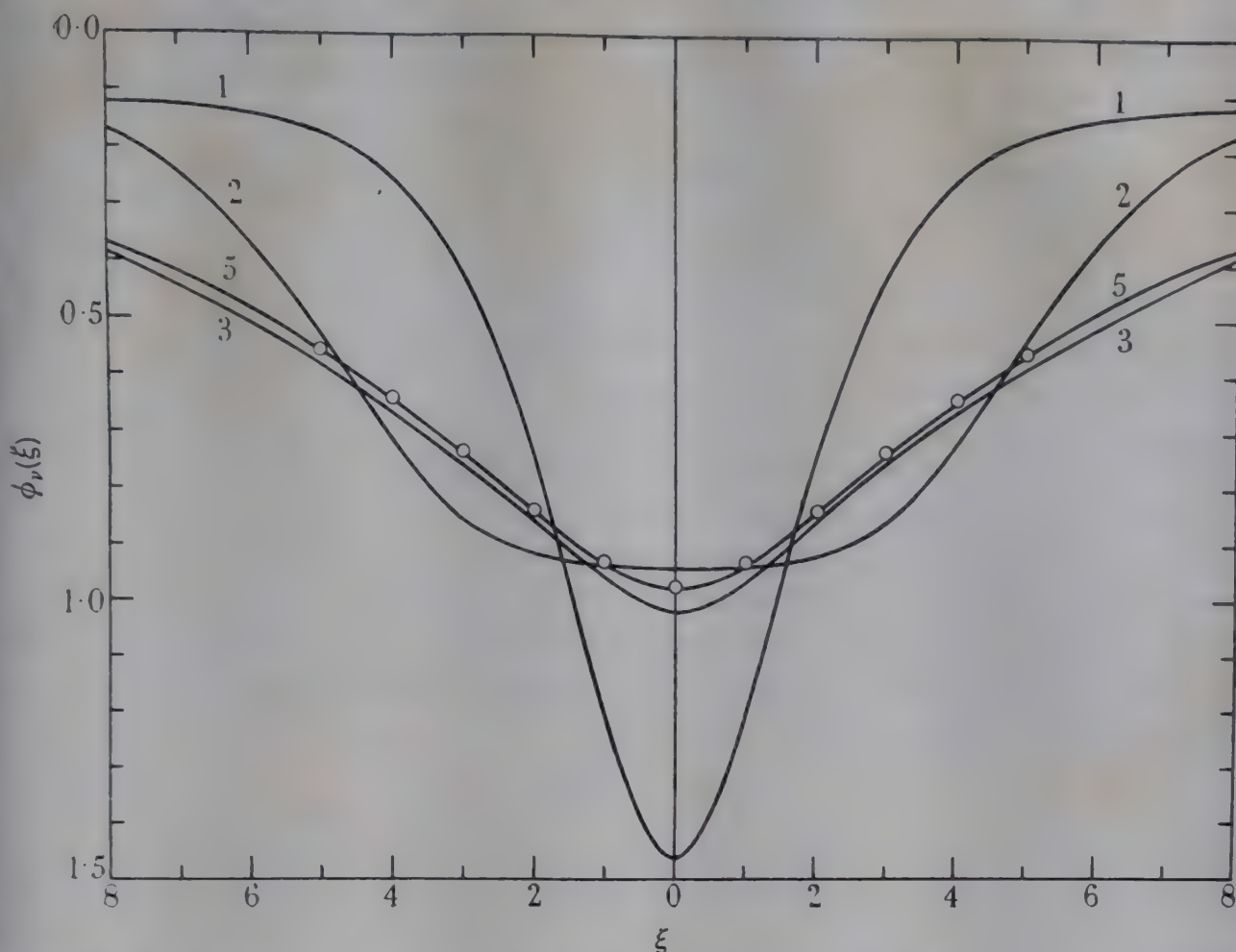


FIGURE 2. Successive approximations to the solution, $\phi(\xi)$, of the transport equation for $y=5$; see equation (3.6). The numbers on the curves denote the stage of approximation, v . The circles show the values of $\phi_0(\xi)$.

Orthogonality condition

A further check on the solutions of the transport equation has been made by a method which is essentially an adaptation of that used by Bloch (1930) to obtain an approximate solution for low temperatures.

Multiplying by $df_0/d\xi$, equation (2.5) may be rewritten as

$$\int_0^y \{G(z) - G(-z)\} dz = -\chi^{\frac{1}{2}} \frac{df_0}{d\xi} + \int_0^y \{H(z) - H(-z)\} dz, \quad (3.9)$$

where

$$G(z) = \frac{z^2 \{\phi(\xi+z) - \phi(\xi)\}}{(e^{-\xi} + 1)(1 - e^{-z})(e^{\xi+z} + 1)}, \quad H(z) = \frac{z^3 \left\{ \frac{z}{y} - \frac{p}{2r} \right\} \phi(\xi+z)}{\chi(e^{-\xi} + 1)(1 - e^{-z})(e^{\xi+z} + 1)}. \quad (3.10)$$

Since (3.9) is satisfied for all values of ξ from $-\eta$ to $+\infty$, the integral of the left-hand side with respect to ξ over this range is equal to the corresponding integral of the right-hand side. Further, since $\eta \gg 1$, and all three terms of (3.9) decrease exponentially for large values of $|\xi|$, the lower limit of the integration with respect to ξ may be replaced by $-\infty$. The integral of the left-hand side is then zero. Consequently the integral of the right-hand side is also zero, that is,

$$\int_{-\infty}^{\infty} \chi^{\frac{1}{2}} \frac{df_0}{d\xi} d\xi = \int_{-\infty}^{\infty} d\xi \int_0^y \{H(z) - H(-z)\} dz. \quad (3.11)$$

This equation represents essentially the orthogonality condition which is satisfied if the inhomogeneous integral equation (3.9) has a finite solution (cf. Brillouin 1931, p. 374; Courant & Hilbert 1924, p. 102).

Substituting for χ from (3.1), noting that $|df_0/d\xi|$ decreases exponentially for large values of $|\xi|$, and that $(ry) \gg 1$, the left-hand side of (3.11) may be evaluated approximately as $-(py)^{\frac{1}{2}}$. Evaluation of the right-hand side in terms of an approximation to ϕ then provides a check on that approximation. Using the final approximations to ϕ it is found in all the cases considered that the right-hand side of (3.11), evaluated by numerical integration, agrees to within 1 % with the left-hand side. The integration over ξ from $-\infty$ to $+\infty$ may be carried out without undue difficulty, since the value of the integral over z decreases rapidly as $|\xi|$ increases.

Low-temperature approximation. The right-hand side of (3.9) decreases as y increases, and to obtain the low-temperature approximation (corresponding to $y \gg 1$) the left-hand side is equated to zero. The solution of this equation is $\phi(\xi) = \text{constant}$, and the value of the constant is determined by (3.11). Assuming χ to be constant at the value corresponding to $\xi = 0$, integration with respect to ξ in (3.11) gives (cf. Wilson 1936, p. 213)

$$\phi(\xi) = -\frac{(p^3 y^5)^{\frac{1}{2}}}{2J_5}, \quad (3.12)$$

where

$$J_5 = \int_0^y \frac{z^5 dz}{(e^z - 1)(1 - e^{-z})}. \quad (3.13)$$

Since the derivation of (3.12) is valid only for $y \gg 1$, the upper limit of the integral in (3.13) may be replaced by ∞ , and J_5 evaluated as $5! \zeta(5)$. With this value of J_5 , equation (3.12) with (2.9) and the Fermi-Dirac approximation gives the low-temperature approximate expression for σ quoted in §2.

4. CALCULATION OF CONDUCTIVITY

The conductivity, σ , may be expressed from (2.9), (3.1) in the reduced form

$$\sigma' = \sigma/\beta = y^{\frac{1}{2}} \int_{-\eta}^{\infty} \left\{1 + \frac{\xi}{ry}\right\}^{\frac{1}{2}} \phi(\xi) \frac{df_0}{d\xi} d\xi, \quad (4.1)$$

where

$$\beta = p^{\frac{1}{2}} \alpha. \quad (4.2)$$

For a given metal β is a constant, and in particular is independent of temperature. From (3.3), $|df_0/d\xi|$ decreases exponentially for $|\xi| \gg 1$; further, $\eta \gg 1$, $(ry) \gg 1$. Consequently, (4.1) may be written, to a sufficient approximation, as

$$\sigma' = y^{\frac{1}{2}} \int_{-\infty}^{\infty} \phi(\xi) \frac{df_0}{d\xi} d\xi + \frac{3y^{\frac{1}{2}}}{2ry} \int_{-\infty}^{\infty} \xi \phi(\xi) \frac{df_0}{d\xi} d\xi. \quad (4.3)$$

Now, as indicated in §3, $\phi(\xi)$ is, to within the required accuracy, an even function of ξ , and $df_0/d\xi$ is also an even function (cf. (3.3)), so (4.3) may be reduced to

$$\sigma' = 2y^{\frac{1}{2}} \int_0^{\infty} \phi(\xi) \frac{df_0}{d\xi} d\xi. \quad (4.4)$$

For convenience in evaluating the integral in (4.4) the range of integration may be divided into two parts: $\xi = 0$ to 5, and $\xi = 5$ to ∞ . The contribution to the integral from the first part may be determined by numerical integration using the final approximations to ϕ . The contribution from the second part is relatively small, and ϕ may be replaced by a constant value, say ϕ_c , limits to which are determined by $\phi_1(5) > \phi_c > \phi(5)$. This part of the integral is then given by

$$\int_5^\infty \phi(\xi) \frac{df_0}{d\xi} d\xi = -\phi_c/(e^5 + 1). \quad (4.5)$$

For small values of $|\xi|$, $|\phi| < |\phi_1|$, and for large values, $|\phi| > |\phi_1|$ (see figure 2). This relation between ϕ_1 and ϕ , which is general, and which applies not only to $y = 5$, suggests, in conjunction with (4.4), that for the smaller values of y , for which ϕ does not differ greatly from ϕ_1 , the value of σ' may be approximately the same when calculated from ϕ_1 as when calculated from ϕ . For $y = 2$ the two values of σ' calculated in this way agree to within 1 %. For $y = 3$ and 5 the values calculated with ϕ_1 are greater than those calculated with ϕ by about 4 and 16 % respectively. It may be concluded that for $y < 2$ it is sufficiently accurate in calculating σ' to replace ϕ by ϕ_1 in (4.4).

Values of σ' have been calculated from (4.4) for $y = 0.5(0.1) 2, 3, 5$, using the values of ϕ and ϕ_1 determined as described in § 3. By interpolation in the σ', y table values of σ' have been determined for $T/\theta_D = 1/y = 0.2(0.1) 1.0, 1.2(0.2) 2.0$. The results are shown in table 2 in the reduced form $\sigma'_1/\sigma'_y = \rho_T/\rho_\theta$, where ρ_T, ρ_θ are the specific electrical resistance at $T^\circ \text{K}$, $\theta_D^\circ \text{K}$, respectively. This reduced form is convenient for comparison with experimental results. Values shown in brackets are somewhat less certain than the others owing to the difficulty of interpolation.

Table 2 also contains values of $(T^5/\rho_\theta)(\rho_T/T^5)_0$ and $(T/\rho_\theta)(\rho_T/T)_\infty$, where $(\rho_T/T^5)_0$ and $(\rho_T/T)_\infty$ are limiting values corresponding to $(T/\theta_D) \rightarrow 0$ and $(T/\theta_D) \rightarrow \infty$, respectively (cf. equations (2.13), (2.12)). It may be noted that although the values of $(T^5/\rho_\theta)(\rho_T/T^5)_0$ and $(T/\rho_\theta)(\rho_T/T)_\infty$ correspond to the low- and high-temperature approximations they could not be obtained from those alone, since neither of the approximations is valid at $T = \theta_D$, and so ρ_θ could not be determined from them.

The values shown in table 2 are considered in relation to the experimental results in § 5. Although values of ρ_T/ρ_θ cannot be determined accurately from the table for $0.2 > (T/\theta_D) > 0.05$, the values given show the temperature variation of resistance over practically the whole of the temperature range of particular interest (cf. § 1, and figures 3 and 4).

The values of $(\rho_T/\rho_\theta)_G$ in table 2 are from the semi-empirical formula suggested by Grüneisen (1930), which may be expressed in the form

$$\left(\frac{\rho_T}{\rho_\theta}\right)_G = \left(\frac{T}{\theta_D}\right)^5 \frac{J_5(\theta_D/T)}{J_5(1)}, \quad (4.6)$$

where $J_5(\theta_D/T)$ is given by (3.13). Tables of values of J_5 are given by Grüneisen (1933). Equation (4.6) may be obtained by assuming that the approximate expression for the conductivity derived from (3.12) is valid over the whole temperature range. The derivation of (3.12) provides no justification for this assumption. Although (4.6)

reduces to the theoretical approximations for $y \gg 1$ and $\ll 1$, for intermediate temperatures it can only be regarded as an empirical interpolation formula. The Grüneisen values agree with the calculated values to within 1 % for $(T/\theta_D) \geq 0.7$ and ≤ 0.05 , but for temperatures between these limits there are significant differences (see table 2).

TABLE 2. SPECIFIC ELECTRICAL RESISTANCE AS A FUNCTION OF TEMPERATURE

ρ_T, ρ_θ , calculated specific resistance at $T^\circ \text{K}$, $\theta_D^\circ \text{K}$;
 $(\rho_T/T^5)_0$, value from the low-temperature approximation (see equation (2.13));
 $(\rho_T/T)_\infty$, value from the high-temperature approximation (see equation (2.12));
 $(\rho_T/\rho_\theta)_G$, values from Grüneisen's (1933) semi-empirical formula (see equation (4.6));
 θ_D , Debye characteristic temperature.

$\frac{T}{\theta_D}$	$\frac{\rho_T}{\rho_\theta}$	$\frac{T^5(\rho_T/T^5)_0}{\rho_\theta}$	$\frac{T(\rho_T/T)_\infty}{\rho_\theta}$	$\left(\frac{\rho_T}{\rho_\theta}\right)_G$
0.05	—	0.000 164 ₆	—	0.000 164 ₃
0.10	—	0.005 27	—	0.004 92
0.15	—	0.040 01	—	0.026 62
0.2	0.060 2	0.168 6	0.212	0.068 0
0.3	(0.169)	—	0.318	0.181 2
0.333	0.210 ₅	—	0.353	0.222 1
0.4	(0.295)	—	0.423 ₅	0.304 5
0.5	0.419	—	0.529	0.426 5
0.6	0.539	—	0.635	0.545 4
0.7	0.658	—	0.741	0.661 7
0.8	0.774	—	0.847	0.775 9
0.9	0.888	—	0.953	0.888 5
1.0	1.000	—	1.05 ₀	1.000
1.2	1.22 ₂	—	1.27 ₁	1.220
1.4	1.44 ₁	—	1.48 ₂	1.438
1.6	1.65 ₇	—	1.69 ₄	1.654
1.8	1.87 ₃	—	1.90 ₆	1.869
2.0	2.08 ₈	—	2.11 ₈	2.084

Range of validity of results

Solutions of the transport equation, (2.5), on which are based the calculation of conductivity, have been obtained for particular values of the parameters p and r , and the limitations which this imposes on the applicability of the results must be considered.

Examination of the derivation shows that the form of the solution of (2.5) does not depend critically upon the value of r , equal to θ_F/θ_D , provided that $r \gg 1$. The Debye temperatures of most metals are such that a small value of r could arise only from a value of θ_F corresponding to electrons distributed in a narrow band, and this case is not treated here. Moreover, an unfilled narrow band is ordinarily superposed on an unfilled wide band, and the relative contribution to the conductivity from the narrow band is then usually very small. The use of the particular value $r = 300$ thus imposes no serious limitation on the range of applicability of the results.

The solution of (2.5) is more sensitive to the value of $p = (2q^2)^{\frac{1}{2}}$, due to the fact that $\chi(2.6)$ is proportional to p . Although no detailed calculations have been made of the

conductivity for values of p other than $2\frac{1}{2}$ (corresponding to one conduction electron per atom), examination of the derivation and of the general run of the results suggests that the temperature variation of conductivity would not be very different from that for $p = 2\frac{1}{2}$.

One further general point may also be noted. Reference has been made only to conduction by electrons, but the results obtained apply equally well to conduction by 'holes' in an otherwise full band, provided that the density of states curve is of standard form at the top of the relevant band.

5. COMPARISON WITH EXPERIMENT AND CONCLUSION

The experimental results to be considered are taken from data given by Grüneisen (1928), Onnes & Tuyn (1929) and Meissner (1935). The results are usually given as the ratio, R_T/R_0 , of the resistance of a specimen at the temperature to the resistance of the same specimen at 0°C . This approximates sufficiently closely to the ratio of the specific resistances, ρ_T/ρ_0 . The values considered here have been 'corrected' by Matthiessen's rule for the temperature-independent 'residual' resistance, and therefore correspond to the 'ideal' resistance. In general the results are for specimens of a high degree of purity and the residual resistance is small compared with the observed resistance except at very low temperatures. Consequently, at intermediate and high temperatures deviations from Matthiessen's rule, such as have been observed by Grüneisen (1933) and considered theoretically by Dube (1938), would have only a slight effect on the ideal resistance.

The experimental and theoretical results may be conveniently compared by using the reduced form ρ_T/ρ_θ . In order to determine this quantity from the given values of ρ_T/ρ_0 , values of ρ_θ/ρ_0 have been determined by graphical interpolation, using values of θ_D from Borelius (1935, p. 252) and Seitz (1940, p. 110). The experimental values of ρ_T/ρ_θ as a function of T/θ_D and the corresponding theoretical curve (cf. table 2) are shown in figure 3. It may be noted that errors in the somewhat uncertain values of θ_D have only a slight effect on the experimental (ρ_T/ρ_θ) , (T/θ_D) relation, since a change in the value adopted for θ_D affects both quantities in approximately the same way.

The most striking feature of figure 3 is the good agreement between the experimental points and the theoretical curve over practically the whole temperature range considered. An alternative method of representation, which brings out more clearly the deviations between the theoretical and experimental results, is suggested by the theoretical linear relation between ρ_T/ρ_θ and T/θ_D at high temperatures. In figure 4 the reduced specific resistance $(\rho_T/\rho_\theta)/(T/\theta_D)$ is shown as a function of T/θ_D .

In figures 3 and 4 the agreement between the theoretical curves and the experimental points for the noble metals, copper, silver and gold, is very good except at high temperatures, where the experimental values lie above the curves. The observed ρ_T , T curves are convex to the T -axis in this region, and this deviation from a linear relation between ρ_T and T has been ascribed by Mott & Jones (1936, p. 268) to a decrease of the effective Debye characteristic temperature with increasing temperature. No accurate theoretical estimate of this effect has been made, but Mott & Jones (1936,

p. 17) give a simplified treatment of the effect of thermal expansion on the characteristic vibration frequencies of a lattice. From this they show (1936, p. 269) that the variation of the characteristic temperature is of the right order of magnitude to account for the observed curvature in the ρ_T, T curves for the noble metals, assuming that $(\rho_T/T) \propto 1/\theta_D^2$ (cf. equation (2.12)).



FIGURE 3. Specific electrical resistance as a function of temperature. Theoretical curve and experimental points. For sources of the experimental values see text. The broken line gives $\frac{T(\rho_T/T)_\infty}{\rho_\theta}$, and the dot-dash curve $\frac{T^5(\rho_T/T^5)_0}{\rho_\theta}$; cf. table 2 and related text. The dotted curve is interpolated. ρ_T, ρ_θ , specific resistance at $T^\circ \text{ K}, \theta_D^\circ \text{ K}$; θ_D , Debye characteristic temperature.

	θ_D		θ_D		θ_D
$\triangle$ Li	400	$\odot$ Cu	315	$\square$ Pb	88
∇ Na	150	$\boxtimes$ Ag	215	$\times$ Pt	225
$\diamond$ K	100	$\otimes$ Au	170	$+$ Pd	275

The agreement between experimental and theoretical results in figures 3 and 4 is not so good for the alkali as for the noble metals. For sodium, in particular, there are significant deviations. It is possible, however, to obtain fairly close agreement with the theoretical curves by adopting a value for θ_D of about 200° K , as compared with the mean value of about 150° K determined from specific heat measurements. This disagreement is surprising, but it is noteworthy that low-temperature specific heat measurements by Pickard & Simon (1948) indicate that, in addition to an anomaly in the specific heat-temperature curve at about 7° K , the effective value of θ_D increases with temperature at least up to 25° K ; the value at 2° K is given as $\theta_D = 88^\circ \text{ K}$, and at 25° K as $\theta_D = 156^\circ \text{ K}$. On the basis of evidence similar to this

Mott & Jones (1936, p. 11) have suggested that for the alkali metals 'the divergence between the true vibration spectrum and the Debye form is particularly great'.

For the transition metals, platinum and palladium, the experimental and theoretical values agree well except at high temperatures, for which the experimental points lie below the theoretical curves. There is strong evidence to indicate that in these metals the electrons outside full bands are distributed in two overlapping bands, the s and d bands. The current is carried mainly by the electrons in the s band, since the effective mass of the 'holes' in the d band is much larger (cf. Mott & Jones 1936,

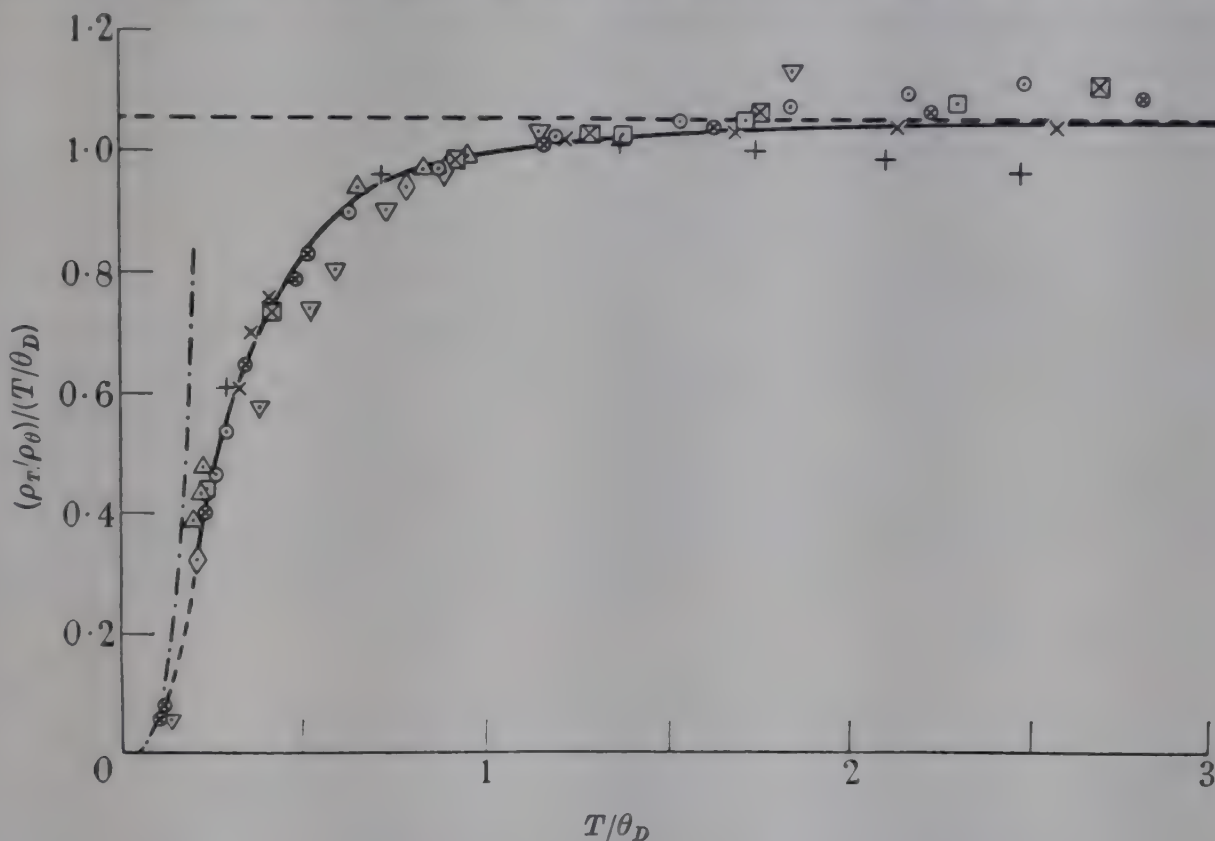


FIGURE 4. Reduced specific electrical resistance as a function of temperature. Theoretical curve and experimental points. For sources of the experimental values see text. The broken, dot-dash and dotted curves correspond to those similarly shown in figure 3. For symbols see figure 3.

p. 267). It has been suggested by Mott (1935, 1936*a, b*) that transitions of electrons from the s band to the d band, due to 'collisions', may play an important part in determining the conductivity of these metals, and it is shown that in this case the estimated ρ_T, T curve is concave to the T -axis, as is observed (cf. Mott & Jones 1936, p. 269). This treatment has, however, been criticized by Wilson (1938), and it cannot be said that a completely satisfactory solution of the problem has yet been attained. Detailed discussion of this particular question is outside the scope of this paper, but two points may be noted. First, it may be shown from (4.1) and the preceding considerations that the temperature variation of resistance of a metal in which the electrons (or 'holes') outside full bands are distributed in two or more bands of standard form is the same as in the case of electrons distributed in a single band, provided that θ_F for each of them is much greater than θ_D , that the numbers of electrons (or 'holes') in the various bands are independent of temperature, and that

there are no transitions between the bands due to 'collisions'. Secondly, in actual metals in which there are overlapping bands there may be a redistribution of electrons among the bands with increasing temperature. This 'transfer effect', which has been considered in detail in relation to the associated magnetic and thermal properties by Wohlfarth (1949), may lead to a relative increase in the conductivity with temperature, and consequently the ρ_T, T curve may be concave to the T -axis.

The principal reason for including the experimental points for lead in figures 3 and 4 is that this metal is representative of a large number of metals for which the temperature variation of electrical resistance agrees well with the theoretical curves of figures 3 and 4, although there is no evidence to suggest that their electronic structures are such as to satisfy the basic assumptions involved in the theoretical treatment; in particular, that the 'conduction' electrons are distributed in a band of standard form (cf. (2.4)). To avoid confusion the experimental points for only a few metals are included in figures 3 and 4 (although even so there is some overlap of the points for different metals). However, since the calculated curves agree fairly well with the Grüneisen curves over most of the temperature range, the comparisons which have been made between experimental values for a large number of metals and the Grüneisen curves (see, for example, Grüneisen 1933; Meissner 1935) serve as a rough comparison between the theoretical curves and experimental values. The good agreement which is found in many cases indicates that in general the temperature variation of electrical conductivity in metals depends only slightly on the electronic structure of the metal, and is determined primarily by the temperature variation of the lattice vibrations.

Concluding remarks

The primary aim in this paper has been to extend the detailed development of the Bloch treatment of the temperature variation of electrical conductivity, so as to cover a much wider temperature range than previously, and to test the applicability of the treatment to actual metals. It will be clear from the preceding discussion that in the main the agreement between calculated and observed values is good over the whole temperature range. Although the agreement is in no case perfect the deviations may in general be ascribed to factors of secondary importance, and it may be concluded that the theoretical treatment covers satisfactorily the main features of the observed temperature variation of conductivity.

In this paper attention has been restricted to the consideration of electrical conductivity, but the Bloch treatment may be readily generalized to deal with more complex transport effects, such as thermal conductivity and thermoelectric power. In these cases integral equations arise analogous to that treated here, and both the methods and the detailed results considered here may be of value in the treatment of these more complex effects.

I am indebted to Professor E. C. Stoner, F.R.S., for his constant guidance and encouragement in this and other work. My thanks are also due to Dr E. P. Wohlfarth for many helpful discussions; to Mr B. A. Lilley for criticism of the draft manuscript; and to the Department of Scientific and Industrial Research for grants.

APPENDIX

A variational treatment of the Bloch equation

Kohler (1949) develops a variational treatment by which successive approximations to the electrical conductivity may be determined from the transport equation without the distribution function being obtained explicitly (see also Kohler 1941, 1948). This method is widely different from that adopted in the present paper, but since both treatments are based on essentially the same initial equation (cf. Kohler 1949, p. 681), they should lead to the same calculated temperature variation of conductivity. Kohler concludes from his variational treatment that the calculated conductivity is, to a close approximation, equal to that given by Grüneisen's empirical formula over the whole temperature range. This result is significantly different from the results of the present paper (cf. table 2).

The fact that the two treatments lead to different results may arise either from inadequacies in one or other of the methods, or from errors in their application. Careful examination of Kohler's paper indicates that there is an algebraical error in the derivation of the expression given for the conductivity (Kohler 1949, pp. 686, 688). This expression may be written in the form

$$\sigma = \sigma_G(1 + \lambda_1 + \lambda_2 + \lambda_3 + \dots), \quad (\text{A.1})$$

where σ_G is the Grüneisen value. It is implied that λ_3 and subsequent terms are negligible, and for λ_1 and λ_2 Kohler obtains

$$\left. \begin{aligned} \lambda_1 &= \left(\frac{\pi^2}{2} - a \right)^2 \left\{ \frac{\pi^2}{3} + 2a - \frac{1}{6} \left(\frac{J_7}{J_5} \right) \right\}^{-1} \gamma^2, \\ \lambda_2 &= \frac{1}{18} \left(\frac{J_7}{J_5} \right) \gamma^2, \end{aligned} \right\} \quad (\text{A.2})$$

where $J_n = \int_0^y \frac{z^n dz}{(e^z - 1)(1 - e^{-z})}$; $a = 2^{-\frac{1}{2}} y^2$; $y = \theta_D/T$; $\gamma = kT/\zeta \doteq T/\theta_F$.

For the relevant temperature range $\gamma \ll 1$ (cf. table 1, §3), so both λ_1 and λ_2 are small for all values of y , and σ is closely equal to σ_G . The derivation of (A.2) is not given in detail. On working through the derivation along the lines indicated by Kohler, the same expression is obtained for λ_1 , but for λ_2 a different expression, namely,

$$\lambda_2 = \left(\frac{J_7}{J_5} \right)^2 \left\{ 24\pi^2 a + \frac{16}{5}\pi^4 + (6a - 2\pi^2) \left(\frac{J_7}{J_5} \right) - \left(\frac{J_7}{J_5} \right)^2 + \frac{3}{10} \left(\frac{J_9}{J_5} \right) \right\}^{-1}, \quad (\text{A.3})$$

is obtained. (Thanks are due to Mr B. A. Lilley for kindly checking the derivation of (A.3).) The numerical value of λ_2 from (A.3) (of order 10^{-2}) is generally much greater than that from Kohler's expression (of order 10^{-7}), since (A.3) does not contain γ as a factor. Values of λ_1 and λ_2 , calculated from (A.2) and (A.3) with values of J_7 and J_9 determined by numerical integration and values of J_5 from Grüneisen (1933), are shown in table 3 for a number of values of y . In evaluating (A.2) the ratio θ_F/θ_D has been taken as 300 (cf. §3).

For comparison with the results of the variational treatment the results obtained by numerical successive approximations (cf. table 2) may be represented in the form

$$\sigma = \sigma_G(1 + L). \quad (\text{A.4})$$

The corresponding values of L are shown in table 3. It is to be noted that the form of (A.4) gives a false impression of the method of derivation of the values in table 2; (A.4) is only introduced to facilitate comparison between the two sets of results.

Since λ_2 is not negligible for a wide range of values of y (cf. table 3) further terms in (A.1) may also have to be taken into account in evaluating σ accurately by the variational method. Moreover, Kohler indicates that all the λ_i are positive, and since L may be written formally as $\Sigma\lambda_i$, it may be concluded from table 3 that the results of the treatment in the present paper are at least not incompatible with those of the variational treatment.

TABLE 3. COMPARISON OF REDUCED CONDUCTIVITIES CALCULATED BY A VARIATIONAL METHOD (THIRD APPROXIMATION) AND BY NUMERICAL SUCCESSIVE APPROXIMATIONS. (FOR EXPLANATION SEE TEXT)

λ_1, λ_2 (A.2), values from (A.2) (cf. (A.1));

λ_2 (A.3), values from (A.3);

L , values calculated from (A.4) and table 2, §4; $y = \theta_D/T$.

y	1	2	3	5
$\lambda_1 \times 10^6$	62	7.0	0.68	0.23
$\lambda_2 \times 10^6$ (A.2)	0.41	0.40	0.39	0.35
$\lambda_2 \times 10^2$ (A.3)	0.1	1.2	3.4	8.1
$L \times 10^2$	<1	1.8	5.5	13

REFERENCES

- Birge, R. T. 1941 *Phys. Soc. Rep. Progr. Phys.* **8**, 90.
 Bloch, F. 1928 *Z. Phys.* **52**, 555.
 Bloch, F. 1930 *Z. Phys.* **59**, 208.
 Borelius, G. 1935 *Handbuch der Metallphysik*, 1/1. Leipzig: Akad. Verlags.
 Brillouin, L. 1931 *Die Quantenstatistik*. Berlin: Springer.
 Courant, R. & Hilbert, D. 1924 *Methoden der mathematischen Physik*, 1. Berlin: Springer.
 Dube, G. P. 1938 *Proc. Camb. Phil. Soc.* **34**, 559.
 Grüneisen, E. 1928 *Handbuch der Physik*, **13**, 1. Berlin: Springer.
 Grüneisen, E. 1930 *Leipziger Vorträge*, p. 46.
 Grüneisen, E. 1933 *Ann. Phys., Lpz.*, **16**, 530.
 Kohler, M. 1941 *Ann. Phys., Lpz.*, **40**, 601.
 Kohler, M. 1948 *Z. Phys.* **124**, 772.
 Kohler, M. 1949 *Z. Phys.* **125**, 679.
 Meissner, W. 1935 *Handbuch der Experimentalphysik*, **11/2**. Leipzig: Akad. Verlags.
 Mott, N. F. 1935 *Proc. Phys. Soc.* **47**, 571.
 Mott, N. F. 1936a *Proc. Roy. Soc. A*, **153**, 699.
 Mott, N. F. 1936b *Proc. Roy. Soc. A*, **156**, 368.
 Mott, N. F. & Jones, H. 1936 *The theory of the properties of metals and alloys*. Oxford: Clarendon Press.
 Nordheim, L. 1931 *Ann. Phys., Lpz.*, **9**, 607.
 Onnes, H. K. & Tuyn, W. 1929 *International Critical Tables*, **6**, 124.
 Pickard, G. L. & Simon, F. E. 1948 *Proc. Phys. Soc.* **61**, 1.
 Seitz, F. 1940 *The modern theory of solids*. New York and London: McGraw-Hill.
 Sommerfeld, A. & Bethe, H. 1933 *Handbuch der Physik*, **24/2**, 333. Berlin: Springer.
 Stoner, E. C. 1939 *Phil. Mag.* **28**, 257.
 Wilson, A. H. 1936 *The theory of metals*. Cambridge University Press.
 Wilson, A. H. 1938 *Proc. Roy. Soc. A*, **167**, 580.
 Wohlfarth, E. P. 1949 *Proc. Roy. Soc. A*, **195**, 434.

Degree of polymerization and chain transfer in methyl methacrylate

BY SADHAN BASU, JYOTIRINDRA NATH SEN AND SANTI R. PALIT

Indian Association for Cultivation of Science, Calcutta, India

(Communicated by M. N. Saha, F.R.S.—Received 27 February 1950)

The chain-transfer constants of twenty-six solvents in the polymerization of methyl methacrylate have been calculated, using Mayo's equation for chain-transfer reaction, from the measured values of degree of polymerization of the product formed in these solvents at different concentrations at 80° C. The increase in transfer constants in hydrocarbons, alcohols, ketones, acids and ester may be effectively accounted for by the increased polarization of the C-H bond at the α -position, by adjacent electron-attracting groups. In the case of chloro-compounds the activity has been found to depend not on the number of halogen atoms in the molecule as was suggested by Mayo but on the extent of C-Cl bond forces.

Further, it is shown that simple chain transfer reactions with solvents cannot fully account for all the observed discrepancies in the kinetic measurements in solution.

Even a decade ago it was assumed that the degree of polymerization is equal to the kinetic chain length; that is, the stabilization of a polymer molecule is accompanied with the destruction of the activity for growth. When, however, the data for kinetic study of polymerization in solution became available, some new facts came to light which necessitated a change in this concept of polymerization kinetics. Thus it was found that the molecular weight of the product formed by the polymerization of styrene in solution was invariably lower than that obtained by polymerization in bulk, although the overall rate was substantially unaffected (Schulz, Dinglinger & Huseman 1939; Schulz & Huseman 1938; Suess, Pilch & Rudorfer 1937; Suess & Springer 1937). Further, according to older concepts, the degree of polymerization should be proportional to the square root of monomer concentration when the catalyst was present. When, however, $\log P$ (P = degree of polymerization) was plotted against $\log A^{\frac{1}{2}}$ (A = initial concentration of monomer, e.g. styrene) a straight line was not obtained (Eyring, Tobolsky, Harman & Hulburt 1943). Further, the average straight line with the various points had a slope definitely less than unity. This indicated that 'some reaction was going on, which terminated the polymer chain without destroying the activity for growth'.

Flory (1937) propounded the chain transfer hypothesis in which it was assumed that the activity of a growing polymer molecule may be transferred to a monomer, polymer or solvent molecule, which in their turn may grow similarly by successive addition of monomer molecule, although the original polymer chain has ceased to grow. Mayo (1943) has used the data of Schulz *et al.* (1938, 1939) and Suess *et al.* (1937) for the kinetics of polymerization of styrene in various solvents in calculating the efficiency of solvents for chain-transfer reaction. Breitenbach & Maschin (1940) have shown that the solvents which are efficient chain-transferring agents very often enter into combination with the polymer formed.

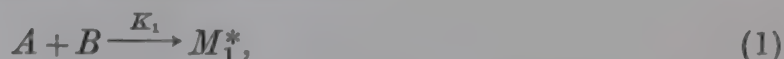
In the present paper the efficiency of various solvents as chain-transferring agents in the uncatalyzed polymerization of methyl methacrylate has been studied and the

results compared with those given by Mayo for the polymerization of styrene in solution. Further, attempts have been made to get a clearer idea as to the chemical nature of the solvent transfer reactions in polymerization of methyl methacrylate.

THEORETICAL

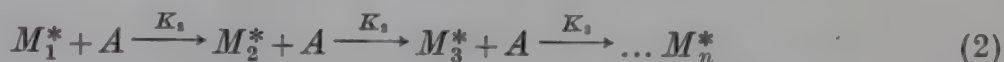
Before proceeding to present the experimental data, it would be advantageous to give a general kinetic consideration of the chain-transfer phenomenon in the polymerization reactions in solutions.

Let us take the case of bimolecular initiation reaction:

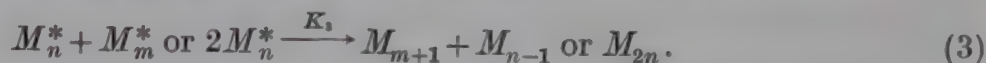


where A , B and M_1^* stand for concentrations of monomer, catalyst and activated molecule respectively. In case of uncatalyzed reaction B is replaced by A .

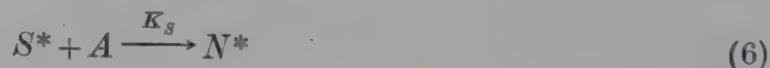
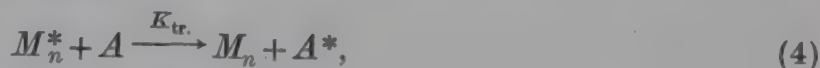
The second step is the propagation of chain by successive addition of monomer A to activated molecule M_1^* :



Ultimately the activity is lost and the growing active molecule is stabilized:



If, however, the active molecule M_n^* be deactivated by transferring the activity to a monomer or solvent molecule, then



where S stands for a concentration of solvent molecule and N^* the active radical formed by the interaction of the active solvent molecule and monomer. Thus we see that though the polymer chain is terminated by the reactions (3), (4) and (5), the kinetic chain is terminated only by the step (3) and the activity is preserved in the steps (4) and (5).

The number average degree of polymerization is equal to the ratio of velocity of propagation to that of termination of the polymer chain, i.e.

$$\begin{aligned} \bar{P} &= \frac{V \text{ propagation}}{V \text{ termination of polymer chain}} \\ &= \frac{V \text{ propagation}}{V \text{ term. (3)} + V \text{ transfer (4), (5)}} \quad (7) \end{aligned}$$

$$= \frac{K_p AC^*}{K_3 C^{*2} + K_{tr.} AC^* + K_{tr.}' C^* S} \quad (8)$$

where C^* is equal to overall concentration of active free radicals. Now when the rate of production of active centres becomes equal to the rate of their destruction, the system reaches a stationary state. Hence at that point

$$K_1 AB = K_3 C^{*2}, \quad (9)$$

or
$$C^* = \left(\frac{K_1}{K_3} AB \right)^{\frac{1}{2}} \quad (10)$$

Substituting the value of C^* from (10) to (8) and inverting we have

$$\frac{1}{\bar{P}} = \frac{(K_1 K_3 [B])^{\frac{1}{2}}}{(K_2 + K_{tr.}) [A]^{\frac{1}{2}}} + \frac{K_{tr.}}{(K_2 + K_{tr.})} + \frac{K_{tr.}'}{(K_2 + K_{tr.})} \frac{[S]}{[A]}, \quad (11)$$

where $[]$ indicates concentrations in g.moles/l. In the case of bulk polymerization $[S]$ and $K_{tr.}'$ are zero, hence we get

$$\frac{1}{\bar{P}_0} = \frac{(K_1 K_3 [B])^{\frac{1}{2}}}{(K_2 + K_{tr.}) [A]^{\frac{1}{2}}} + \frac{K_{tr.}}{K_2 + K_{tr.}} \quad (12)$$

For uncatalyzed bulk polymerization $B = A$, and we have

$$\frac{1}{\bar{P}_0} = \frac{(K_1 K_3)^{\frac{1}{2}}}{K_2 + K_{tr.}} + \frac{K_{tr.}}{K_2 + K_{tr.}} \quad (13)$$

The equation for uncatalyzed polymerization in solution reduces to

$$\frac{1}{\bar{P}} = \frac{K_{tr.}' [S]}{(K_2 + K_{tr.}) [A]} + \frac{1}{\bar{P}_0}, \quad (14)$$

or
$$\frac{1}{\bar{P}} = C \frac{[S]}{[A]} + \frac{1}{\bar{P}_0} \quad (15)$$

where C is a constant and is defined as the chain-transfer constant. This equation in a generalized form (11) has been deduced by Eyring, and the special form (15) has been used by Mayo in the case of styrene. It should be remembered that this equation should apply only at a very low percentage conversion (below 10 %), where monomer concentration may be taken as practically constant so that the initial $[S]/[A]$ value may be used.

If the equation (14) be examined more carefully it becomes clear that intercept of the line on $1/\bar{P}$ axis will be approximately equal to $1/\bar{P}_0$, i.e. the reciprocal of the degree of polymerization for bulk reaction provided the propagation rate after chain transfer remains the same as before it and there is no polymer transfer. If one solvent be more effective as chain-transferring agent than another, other terms remaining constant, $K_{tr.}'$, and consequently C , will be greater in one case than in the other. Hence the value of C is a measure of chain-transferring capacity of the solvent used.

If the activated solvent molecule formed by transfer reaction be unable to initiate further chains, then the reaction (6) will be absent and the activated solvent molecules will destroy each other by mutual combinations. In that case active centres will be lost by reaction (5), and as a result the overall rate will be lower than that where the activated solvent molecules can initiate further chain.

EXPERIMENTAL

Preparation of the monomer

Monomeric methyl methacrylate was prepared by pyrolytic decomposition of Plexiglass, twice distilling the product and collecting the fraction having a boiling range of 99 to 102°C; it was preserved with hydroquinone. Just before working, the stabilized monomer was washed free of hydroquinone with 5 % NaOH solution, then thoroughly washed with water dried over calcium chloride and distilled twice, collecting the fraction having a boiling range of 100 to 101°C.

Purification of solvents used

The solvents used were mostly A.R. and some L.R. quality samples. They were all scrupulously purified by the usual methods, dried and fractionally distilled before use.

Polymerization experiments

Polymerizations were carried out in Pyrex glass ampoules of 20 ml. capacity. These were thoroughly cleansed with chromic acid, washed with water, then with sulphurous acid solution, again with water and dried. Different mixtures (10 ml.) of solvent and monomer were taken in the ampoules, frozen in liquid oxygen, evacuated and sealed. These tubes were then suspended in an oil thermostat at 80°C (accurate within $\pm 0.1^\circ\text{C}$). By trial, the time required for each solvent-monomer mixtures for about 10 % conversion was first found. Then a number of tubes with different mixtures of solvent and monomer were suspended in the oil bath. After specified times different tubes were taken out and cooled, the ampoules were broken open and the polymer freed from monomer by double precipitation from benzene solution with methyl alcohol. The polymer was then dried *in vacuo* at 60°C for 2 hr.

Determination of degree of polymerization

The degree of polymerization was determined from the relation $n = k[\eta]^\alpha$, where n = degree of polymerization, $[\eta]$ the intrinsic viscosity, k and α are constants, equal respectively to 2.81×10^3 and 1.32 in benzene solution (Baxendale, Bywater & Evans 1946).

Intrinsic viscosity, $[\eta]$, was obtained by extrapolating the graph η_{sp}/C against C to zero concentration. ($\eta_{sp} = \eta/\eta_0 - 1$, where η is the viscosity of solution and η_0 that of the solvent and C is concentration expressed in g./100 ml. of solvent.)

The viscosity measurement were done in benzene solution at 35°C ($\pm 0.01^\circ\text{C}$), using two Ostwald capillary viscometers with flow times of 189 and 268 sec. respectively with benzene.

RESULTS

Table 1 summarizes the data for the percentage polymerization of methyl methacrylate at 80°C in a number of solvents after 12 hr. The results show fairly convincingly that the overall rate remains fairly constant in all the solvents mentioned below.

In table 2 are given the values of intrinsic viscosity, $[S]/[A]$, $\bar{P}$ (degree of polymerization), $1/\bar{P}$ and C , the chain-transfer constant for various solvents (total being 26) at 80°C in the polymerization of methyl methacrylate.

TABLE 1. EXTENT OF POLYMERIZATION OF METHYL METHACRYLATE IN SOLUTION WITH 50 % MOLE CONCENTRATION OF MONOMER

solvents	percentage polymerized after 12 hr.	solvents	percentage polymerized after 12 hr.
benzene	9.86	methyl ethyl ketone	9.83
ethyl benzene	9.77	carbon tetrachloride	9.74
chlorobenzene	9.87	toluene	9.77
chloroform	9.74		

TABLE 2. CHAIN TRANSFER CONSTANTS FOR METHYL METHACRYLATE POLYMERIZATION IN VARIOUS SOLVENTS AT 80° C

solvents	[S]/[A]	[η]	$\bar{P}$	$1/\bar{P} \times 10^4$	$1/\bar{P}_0 \times 10^4$	$C \times 10^5$
benzene	9.00	2.40	8925	1.12	0.40	0.75
	2.30	3.38	14030	0.71		
	1.00	4.40	19860	0.50		
cyclohexane	9.00	2.20	9545	1.04	0.15	1.00
	4.00	4.00	17520	0.50		
	2.30	4.50	20460	0.48		
methyl-cyclohexane	7.65	1.22	3560	2.80	0.45	1.95
	1.98	2.60	9886	1.00		
	0.85	3.90	16890	0.59		
toluene	4.55	0.95	3783	2.60	0.25	5.25
	2.33	1.85	6330	1.50		
	1.00	3.40	14140	0.70		
ethyl benzene	8.75	0.37	740	13.50	0.80	13.50
	3.40	0.70	1755	5.60		
	2.10	1.00	2807	3.50		
<i>iso</i> -propyl benzene	6.93	0.46	1006	13.70	0.60	19.00
	3.08	0.65	1586	6.30		
	0.95	1.35	4161	2.40		
<i>t</i> -butyl benzene	6.13	1.42	4470	2.20	0.50	2.60
	1.59	2.89	11360	0.88		
	0.68	3.40	14080	0.71		
<i>n</i> -butyl alcohol	10.41	1.10	3186	3.14	0.50	2.50
	2.70	2.40	8925	1.12		
	1.16	2.90	11460	0.88		
<i>iso</i> -butyl alcohol	10.28	1.08	3090	3.24	0.50	2.50
	4.57	1.75	5882	1.70		
	2.67	2.25	8196	1.22		
<i>sec</i> -butyl alcohol	10.38	0.48	1052	9.52	0.80	8.50
	2.69	1.20	3575	2.80		
	1.15	1.63	5334	1.88		
<i>t</i> -butyl alcohol	10.14	1.78	5992	1.67	0.80	1.00
	4.51	2.10	7482	1.34		
	2.63	2.63	10040	1.00		
acetone	7.36	1.41	4407	2.20	0.60	2.25
	4.28	2.09	7406	1.30		
	1.84	2.62	9933	1.00		
methyl ethyl ketone	10.53	0.55	1276	7.80	1.20	7.00
	4.68	0.90	2444	4.00		
	1.17	2.30	4054	2.00		

TABLE 2 (cont.)

solvents	$[S]/[A]$	$[\eta]$	$\bar{P}$	$1/\bar{P} \times 10^4$	$1/\bar{P}_0 \times 10^4$	$C \times 10^5$
methyl- <i>iso</i> -butyl ketone	7.65	0.70	1755	5.60	0.20	7.00
	3.40	1.30	3959	2.50		
	0.85	2.99	11890	0.80		
diethyl ketone	9.18	0.29	546	18.30	1.80	17.75
	2.38	0.66	1614	6.10		
	1.00	0.99	2764	3.60		
acetic acid	16.20	1.10	2475	4.04	0.70	2.40
	7.20	1.52	4883	2.04		
	1.80	2.40	8913	1.12		
ethyl-acetate	9.54	0.92	2500	4.00	1.10	2.40
	4.24	1.14	3333	3.00		
	1.06	1.70	4177	2.39		
<i>iso</i> -butyric acid	9.99	0.45	979	10.21	1.00	9.00
	4.44	0.77	1989	5.02		
	1.11	1.47	4672	2.14		
dioxan	9.00	1.30	3972	2.50	0.30	2.22
	4.00	2.10	7489	1.30		
	2.30	2.90	10300	0.90		
	1.00	4.52	20580	0.40		
butyl chloride	3.75	0.81	2120	4.70	0.20	12.00
	2.70	0.99	2764	3.60		
	1.02	2.20	7929	1.20		
chloroform	6.80	0.46	1010	9.90	0.50	14.00
	3.00	0.85	2268	4.40		
	0.76	2.55	7164	1.30		
carbon tetrachloride	9.00	0.25	451	22.10	0.50	23.93
	1.25	1.00	2810	3.50		
	0.40	2.00	7015	1.40		
propylene chloride	4.00	0.90	2445	4.00	0.20	6.75
	2.45	1.20	3575	2.70		
	1.00	2.32	8535	1.10		
methyl chloroform	9.18	0.72	1819	5.49	0.20	6.00
	6.50	1.10	2475	4.04		
	3.60	1.50	4793	2.08		
tetrachloroethane	9.09	1.40	4366	2.20	0.20	2.00
	4.04	2.90	11430	0.87		
	1.01	4.25	18880	0.52		
chlorobenzene	6.80	1.50	4788	2.08	0.80	2.00
	3.00	1.90	6105	1.43		
	1.80	2.20	9786	1.02		

It is to be noted, however (figures 1 to 5), that although the plot of $1/\bar{P}$ against $[S]/[A]$ gives fairly good straight lines with all the solvents, as is to be expected if the assumption of the solvent transfer reaction be justified, the lines do not pass through the same point for all the cases, not to speak of through the point $1/\bar{P}_0$ for uncatalyzed bulk polymerization (0.4×10^{-4}). There may be several causes for this sort of behaviour. First, we have assumed that the free radical generated from the solvent molecule propagates the chain at the same rate as the original chain, which may not be the case. Secondly, the reaction may take place between the free radical generated

from the solvent with a polymer molecule, the importance of which remains still to be determined. Thirdly, we have used weight-average molecular weight in the calculation of the degree of polymerization, while theoretically the number-average value should be used. But since the spread of the intercepts along $1/\bar{P}$ axis depends on the nature of the solvents, and not on the absolute value of chain-transfer constant, they may very well be attributed to some side reactions other than the chain-transfer reaction with the solvent.

A comparison of the values of chain-transfer constants of various solvents (common to both) in the polymerization of styrene (Mayo 1943) and methyl methacrylate shows that the sequence remains practically the same in both cases, although the absolute values differ appreciably in degree from solvent to solvent.

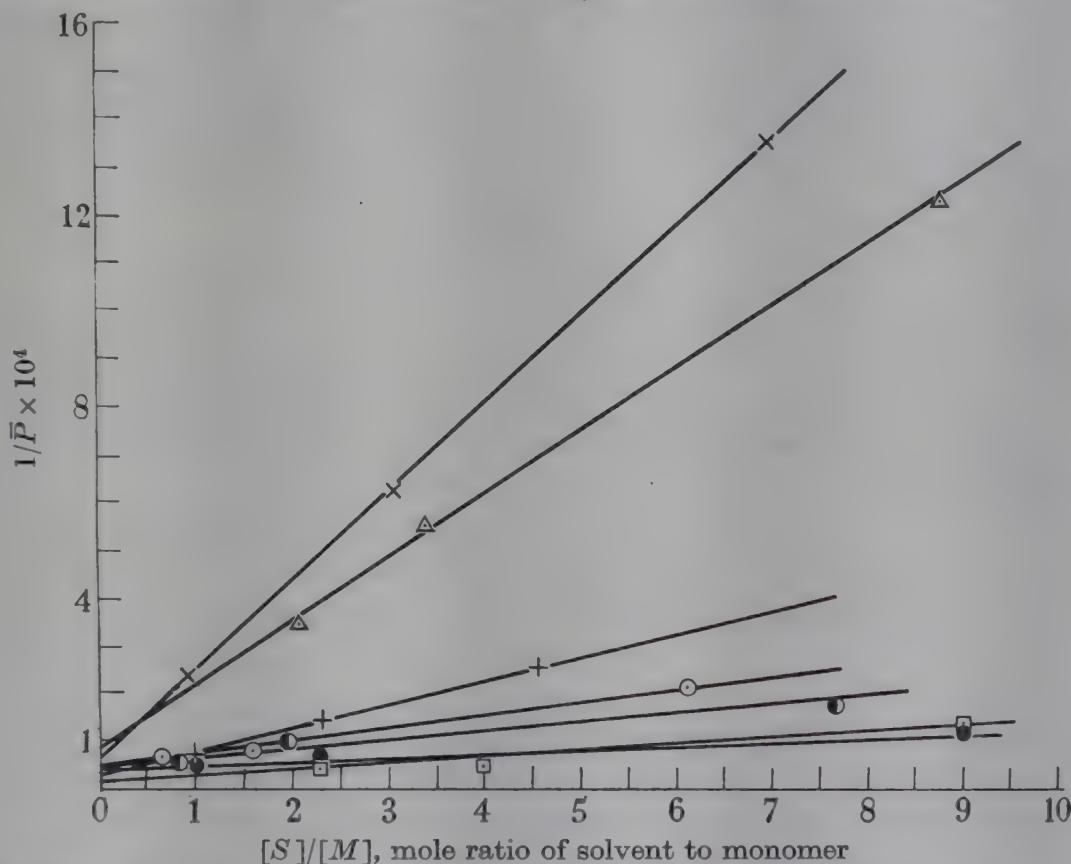


FIGURE 1. $\times$, $\text{C}_6\text{H}_5\cdot\text{CH}(\text{CH}_3)_2$; $\triangle$, $\text{C}_6\text{H}_5\text{CH}_2\cdot\text{CH}_3$; $+$, $\text{C}_6\text{H}_5\cdot\text{CH}_3$; $\odot$, $\text{C}_6\text{H}_5\cdot\text{C}(\text{CH}_3)_3$; $\bullet$, $\text{C}_6\text{H}_{11}\cdot\text{CH}_3$; $\square$, C_6H_{12} ; $\bullet$, C_6H_6

TABLE 3. COMPARISON OF TRANSFER CONSTANTS OF STYRENE WITH THAT OF METHYL METHACRYLATE

solvents	chain-transfer constant $\times 10^5$		
	styrene at 100° C		methyl methacrylate at 80° C
	Schulz <i>et al.</i> (1939)	Suess <i>et al.</i> (1937)	present authors
carbon tetrachloride	—	110.0	23.93
ethyl benzene	14.0	23.0	13.50
toluene	5.3	7.2	5.25
chlorobenzene	—	5.4	2.00
benzene	3.1	4.2	0.75
cyclohexane	3.1	—	1.00

DISCUSSION

Since it is assumed that the solvent molecules do not contribute in any way to the initiation reaction, their sole effect may be presumed to be due to their activity as chain-transferring or chain-terminating agents. Since the rate of polymerization in various solvents (table 1) is not markedly affected, the effect of solvent on the degree of polymerization may be attributed to the chain-transfer reaction as discussed above. Relative efficiency of the various solvents will, therefore, be reflected in their respective chain-transfer constant.

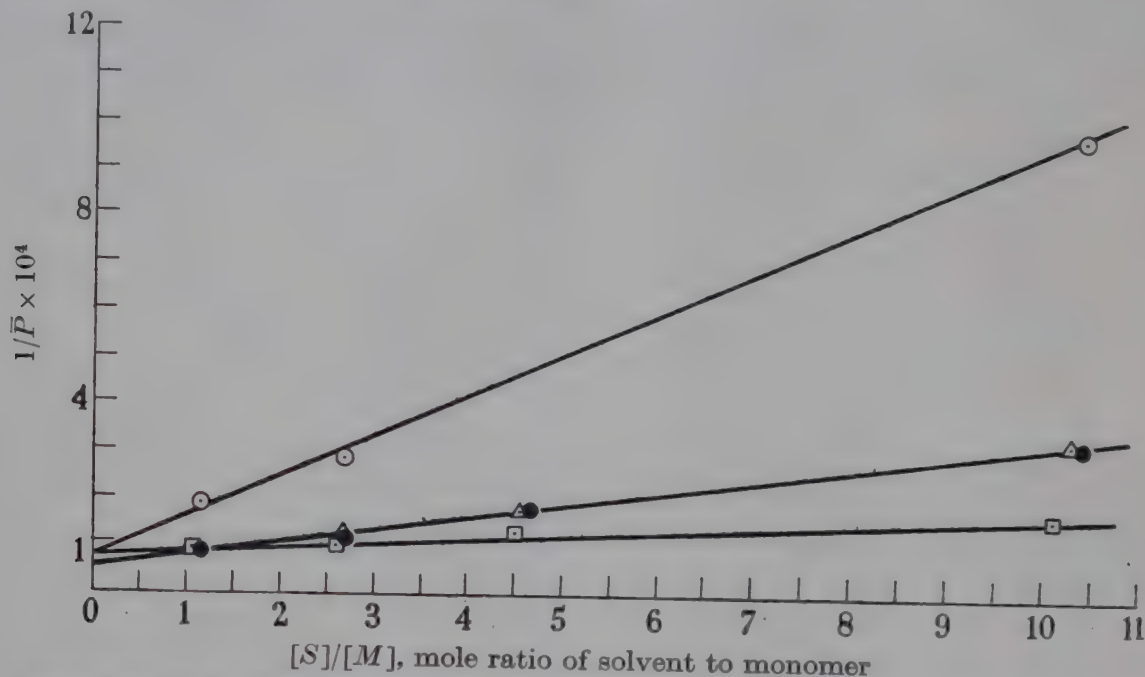


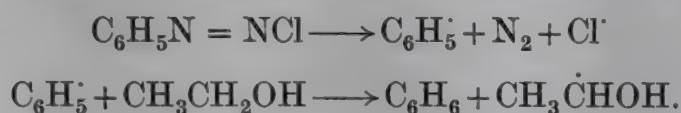
FIGURE 2. O, $\text{CH}_3\text{CH}_2\text{CH}(\text{OH})\text{CH}_3$; Δ , $(\text{CH}_3)_2\text{CHCH}_2\text{OH}$; $\bullet$, $\text{CH}_3\text{CH}_2\text{CH}_2\text{CH}_2\text{OH}$; $\square$, $(\text{CH}_3)_3\text{C.OH}$.

From the chain-transfer constants for various solvents given in the last column of table 2, it would be evident that some solvents are efficient chain-transferring agents and reduce the degree of polymerization enormously, while the others are but feebly effective. The question then arises as to what this transferring activity is due. Price & Tate (1943) have shown that in inhibition, retardation and moderation of polymerization reaction by foreign organic substances, the fragments of the added molecule very often enter into the constitution of the polymer formed. Taking the cue from these observations Mayo (1943) assumed that in the case of hydrocarbons the activity in chain-transfer reaction is due to transfer of an active atom from the solvent molecule to the growing polymer radical, itself at the same time becoming a free radical and propagating the chain further. If this assumption be justified then it may be argued that activity will increase with increase in the mobility of hydrogen atoms in a compound, i.e. with the increased polarization of C-H bond of the molecule by a negative electron attracting group or groups of atoms.

In benzene and cyclohexane, the transfer constants are low, 0.75×10^{-5} and 1×10^{-5} respectively, evidently due to the absence of any active hydrogen atoms. In toluene, however, the side-chain hydrogen atoms are very reactive owing to their attachment to a carbon atom attached directly to the electron-attracting phenyl

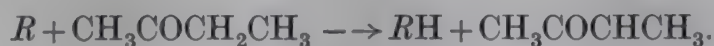
nucleus. As a result the side-chain hydrogen atoms can readily take part in the transfer reaction, and consequently the transfer constant of toluene is high, 5.25×10^{-5} . In methyl cyclohexane this activity is rather low, 1.95×10^{-5} , since the cyclohexane ring is very much less electron-attracting than the phenyl ring, and thus the C-H bonds of the CH_3 groups are but feebly polarized. In ethyl benzene and *iso*-propyl benzene the increased reactivity, 13.5×10^{-5} and 19×10^{-5} respectively, is due to secondary and tertiary hydrogen, the latter being more effective than the former, since it is the experience of the organic chemists that the reactivity of the hydrogen attached to the α -carbon atom increases in the order normal < secondary < tertiary (Glazebrook & Pearson 1936). In the case of tertiary butyl benzene, which has no α -hydrogen, the transfer constant falls back to a very low value 2.6×10^{-5} .

In alcohols, e.g. butyl alcohol series, the active secondary or tertiary hydrogen attached to the carbon carrying the hydroxyl group is mainly effective in chain-transfer reaction. It has been shown by Hey & Waters (1937) that in the decomposition of diazonium chloride in ethyl alcohol, the following reaction takes place:



In the case of polymerization it may be expected that by similar reaction with polymer radical, the alcohols will become effective chain-transferring agents. In *n*-butyl and *iso*-butyl alcohol the reactivity is the same, 2.5×10^{-5} , since both of them contain the same number of active secondary hydrogen atoms attached to the carbon atom carrying the hydroxyl group. In *sec.*-butyl alcohol, however, the highly reactive tertiary hydrogen increases the transfer constant considerably, 8.5×10^{-5} , while in tertiary butyl alcohol the absence of any active hydrogen causes the transfer constants to fall back to a very low value, 1×10^{-5} .

In ketonic compounds the hydrogen atoms at α , α' positions to the carbonyl group are reactive and most probably take part in chain-transfer reaction. It has been shown by Moore & Taylor (1940) that methyl ethyl ketone is attacked by free radical as follows:



In polymerization the chain transfer may take place by a similar mechanism. In methyl ethyl ketone, $\text{CH}_3\text{COCH}_2\text{CH}_3$, the polarization of C-H bonds of the CH_2 group is much greater than those of the CH_3 group in acetone, hence the activity of the former (7×10^{-5}) is greater than that of the latter (2.25×10^{-5}). In methyl *iso*-butyl ketone, $\text{CH}_3\text{COCH}_2\text{CH}(\text{CH}_3)_2$, since the number of active hydrogens is equal to that of methyl ethyl ketone, the transfer constant remains the same (viz. 7×10^{-5}). In diethyl ketone, $\text{CH}_3\text{CH}_2\text{COCH}_2\text{CH}_3$, there are four active CH_2 -hydrogen atoms at α , α' positions, and so the activity is much higher, 17.75×10^{-5} , than those of the other ketones mentioned.

In the case of ketones, however, the spread of the intercepts along the $1/\bar{P}$ axis is much greater than those of the other solvents mentioned above, which indicates that some side-reaction was going on along with the chain-transfer reaction in the former, which has not been accounted for.

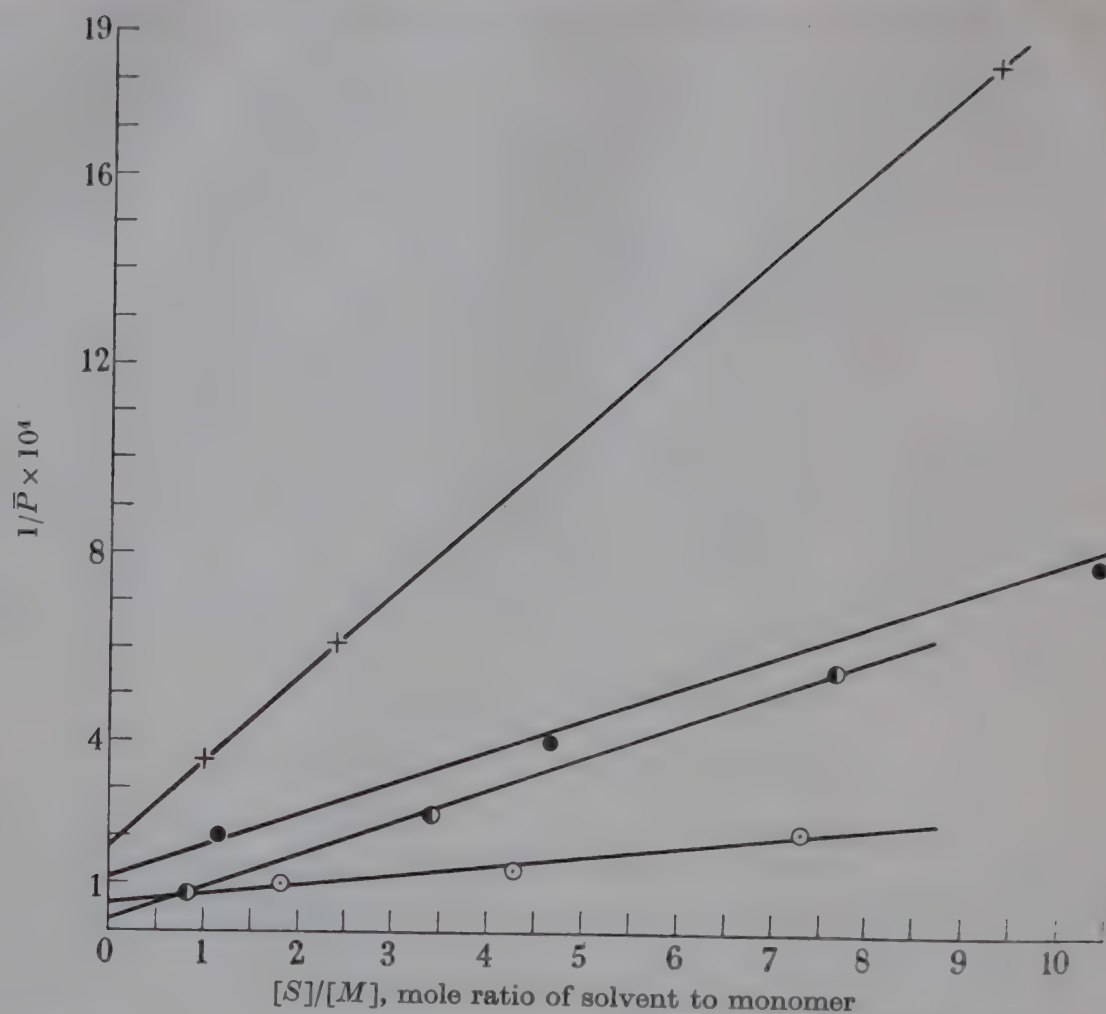


FIGURE 3. $+$, CH₃CH₂COCH₂CH₃; $\bullet$, CH₃COCH₂CH₃; $\odot$, CH₃COCH₂CH(CH₃)₂; $\ominus$, CH₃COCH₃.

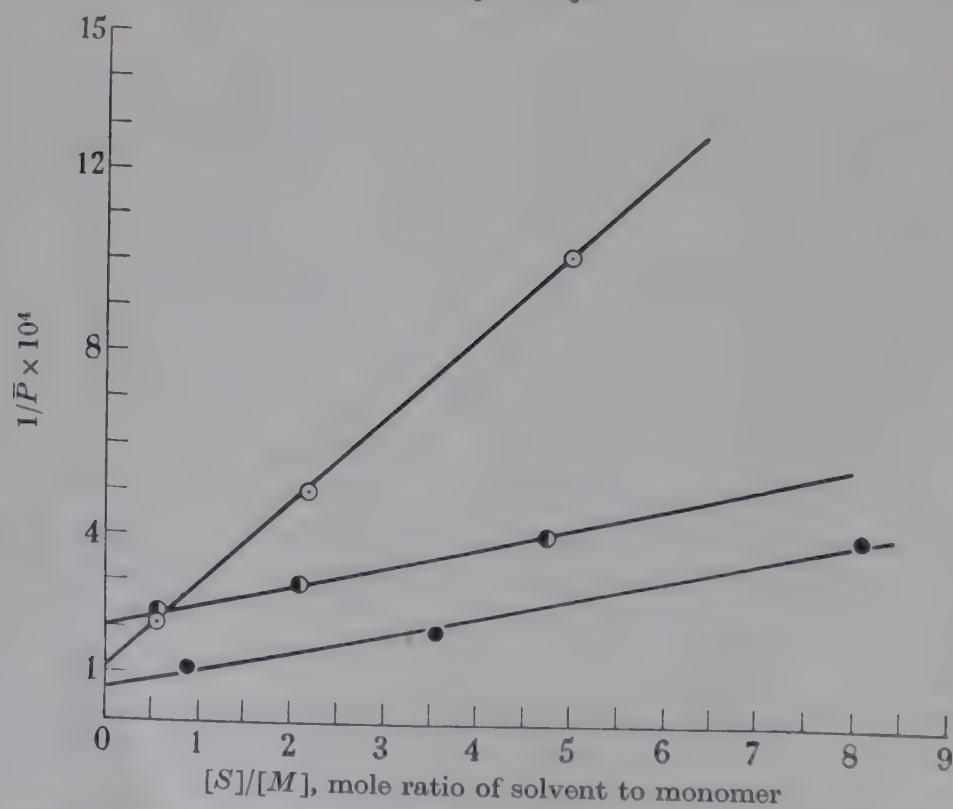


FIGURE 4. $\ominus$, (CH₃)₂CH.COOH; $\odot$, CH₃COOC₂H₅; $\bullet$, CH₃COOH.

The chain-transferring capacity of the acids and esters may also be explained similarly by the active hydrogen hypothesis. In the case of acids the reactivity of the hydrogen increases in the order primary < secondary < tertiary. Thus, *iso*-butyric acid with a tertiary hydrogen is a much more active chain-transferring agent than acetic acid containing only primary α -hydrogen, the transfer constants being

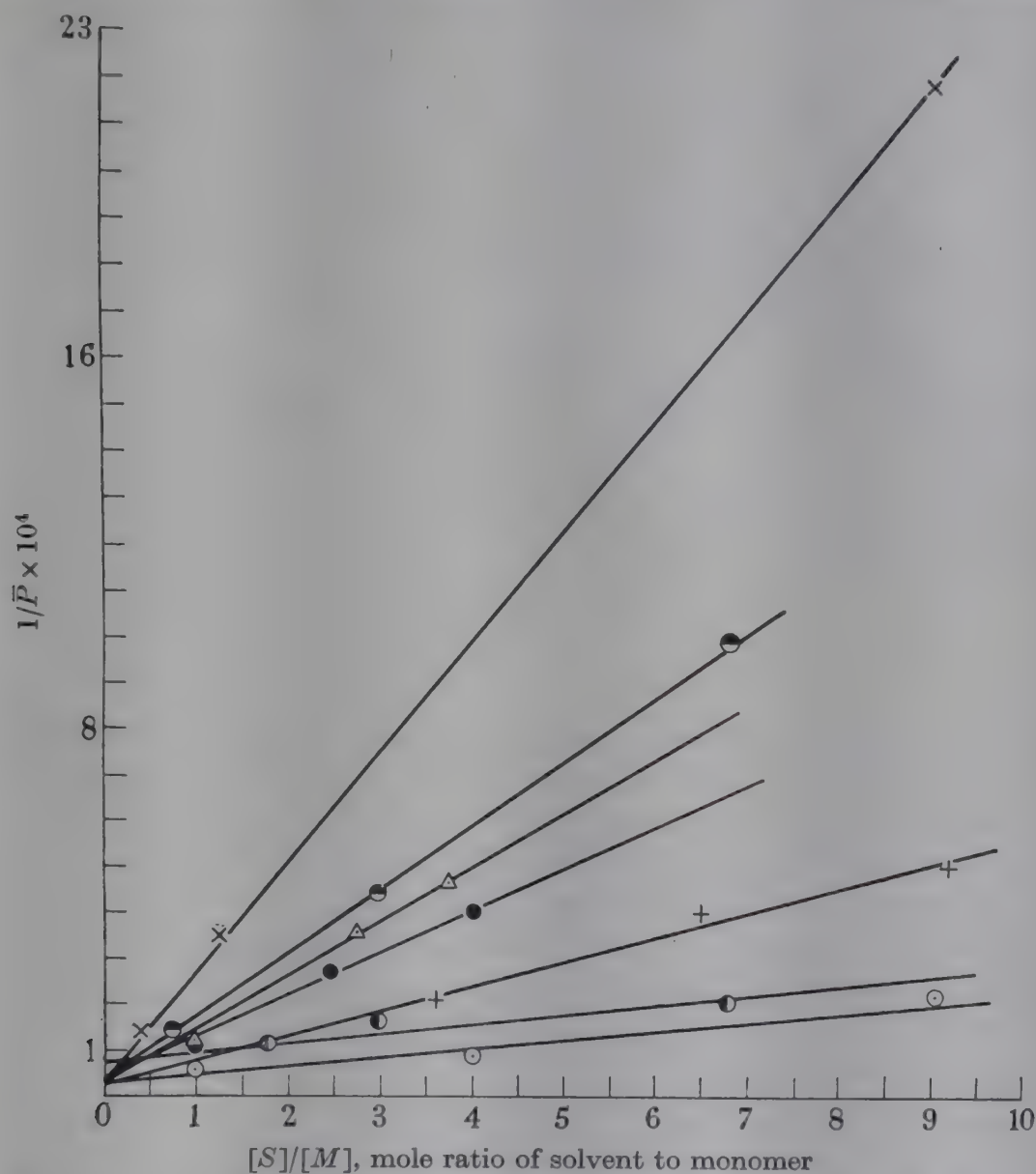


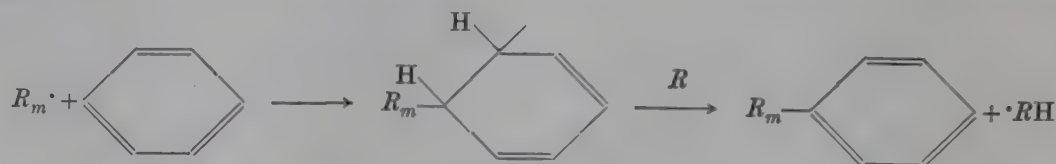
FIGURE 5. $\times$, CCl_4 ; $\bullet$, CHCl_3 ; $\triangle$, $\text{CH}_3\text{CH}_2\text{CH}_2\text{CH}_2\text{Cl}$; $\bullet$, $\text{CH}_3\text{CHCl}.\text{CH}_2\text{Cl}$; $+$, CH_3CCl_3 ; $\bullet$, $\text{C}_6\text{H}_5\text{Cl}$; $\circ$, $\text{CHCl}_2\text{CHCl}_2$.

9×10^{-5} and 2.4×10^{-5} respectively. The α -hydrogen reactivity of acids to free radical attack has been demonstrated in a striking manner by Kharasch & Gladstone (1943) in the formation of succinic acid and its derivatives from aliphatic carboxy acids in presence of acetyl peroxide.

The chain-transfer constants of ethyl acetate, 2.4×10^{-5} , is equal to that of acetic acid, which shows that the α -hydrogens are equally reactive in both cases. It has been shown by Kharasch & Brown (1940), and Price & Schwarcz (1940) also that the reactivity of α -hydrogen of an acid to free radical attack is not affected to any appreciable extent by its conversion to ester. However, in the case of acetic acid and

ethyl acetate, although the chain-transfer constant remains the same, the spread of the intercepts along the $1/\bar{P}$ axis is very great. This may also be due to some other side-reaction apart from transfer reaction.

From the above observations it may be evident that although the high chain-transferring capacities may be attributed to the presence of an active group, or groups, of atoms in the solvent molecules which can be transferred to the growing polymer molecules, the chain-transfer reactions in general may not be exclusively due to them, since even chemically inert molecules like benzene, methyl cyclohexane, *t*-butyl benzene, dioxane, etc., can enter into transfer reactions appreciably. These are most probably due to the peculiarity of the free radical reactions themselves, which do not always conform to the reactivity of a compound from the structural point of view. Price (1946) has attributed the reactivity of benzene to the following reaction:



But it may also be due to reaction of the type



which benzene is known to undergo in presence of lead tetra-acetate (Fieser & Oxford 1942). It is, however, not possible kinetically to decide which one of the mechanism is actually operating since the overall kinetic effect is the same in both cases. Similar is the case with *t*-butyl benzene. The somewhat higher reactivity of methyl cyclohexane is most probably due to α -methylenic hydrogen, which is known to undergo chain reaction with methyl and phenyl radical (Farmer 1942). Price (1946) has shown that even ethers, e.g. $(\text{CH}_3)_2\text{CHOR}$, can undergo the following reaction with free radical:



The reactivity of dioxane may be due to some such reaction. However, these reactions are extremely slow, as evidenced from the low chain-transfer constants of these solvents.

This aspect of the problem of reactivity of inert molecules to free radical attack is very strikingly illustrated by the behaviour of aliphatic chloro compounds in the chain-transfer reaction. Carbon tetrachloride, a chemically inert compound, is a very effective chain-transferring agent, and its activity has been attributed to its ready decomposition under free radical attack. This suggestion has been convincingly supported by the fact that a large variety of compounds which can generate free radicals by decomposition invariably react with carbon tetrachloride and aliphatic chloro-compounds. Thus, it has been shown by Hey & Waters (1937) that a free radical generated by the decomposition of benzoyl peroxide can react with carbon tetrachloride to give chlorobenzene and hexachloro-ethane. In the case of polymerization reactions the carbon tetrachloride molecules are also attacked by

the growing polymer radical leading to a sequence of chain reactions, the only difference in this case is that the CCl_3 radical generated instead of being destroyed by mutual combination, can propagate a new reaction chain. This mechanism has also been supported by Breitenbach & Maschin (1940), and by Mayo (1948), who detected the presence of $\text{Cl}_3\text{C}(\text{CH}_2\text{CHC}_6\text{H}_5)_n\text{Cl}$ molecules in the polymerized product of styrene in carbon tetrachloride, by analytical means.

Seven chlorinated hydrocarbons have been studied in the present case, and it will be observed from figure 5 and table 2 that they differ widely among themselves in their chain-transfer capacity, the order being $\text{CCl}_4 > \text{CHCl}_3 > \text{butyl chloride} > \text{propylene chloride} > \text{CH}_3\text{CCl}_3 > \text{tetrachloroethene}$ and chloro-benzene. It is remarkable that even with such wide divergence in the carbon and chlorine content all the aliphatic chloro-compounds obey the chain-transfer equation (15) with a fair degree of accuracy, thereby justifying the correctness of the assumptions involved. Mayo (1943) assumed that the transfer reactivity of chloro-compounds increases with the increase in the number of chlorine atoms in the molecule, and this apparently holds good in the present case as shown by the order $\text{CCl}_4 > \text{CHCl}_3 > \text{butyl chloride}$. This, however, is not at all true if the chlorine atoms are on different carbon atoms as it would be observed that propylene chloride and tetrachloroethane have chain-transfer constant lower than that of butyl chloride. In fact, tetrachloro-ethane is practically as resistant to free radical attack as chloro-benzene. Further, the ready susceptibility to free radical attack of the hydrogen in chloroform has been surmised by Price (1946) and this is seen here to be corroborated from the results that the replacement of the hydrogen atom in chloroform by methyl group brings about a pronounced decrease of the transfer constant from 14×10^{-5} to 6×10^{-5} .

Thus we see that in the case of chloro-compounds it is not possible to predict from simple considerations the activity of a compound in chain-transfer reaction either by increased chlorine content or by any structural considerations as has been the case with active hydrogen hypothesis. In all probability the energy of C-Cl bonds in different compounds determine their susceptibility to free radical attack, as has been suggested by Dewar (1949). This also becomes evident from the values of force constants of C-Cl bond in different compounds, obtained from Raman spectra, which are in the order $\text{CCl}_4 > \text{CHCl}_3 > \text{propylene dichloride}$ (Remick 1949).

From all these observations it becomes clear that although chain-transfer reaction is a reality and plays a very prominent role in the polymerization reaction in solution, it cannot account for all the peculiarities and discrepancies observed in the kinetic study of the polymerization reactions in solution, as was formerly claimed. The chemistry of free radical reaction, the extent, intensity and kinetics of polymer and monomer transfer reactions in different solvents should be clearly understood before a fruitful attempt could be made to elucidate the polymerization kinetics in solution.

REFERENCES

- Breitenbach, J. W. & Maschin, A. 1940 *Z. phys. Chem. A*, **187**, 175.
 Baxendale, J. H., Bywater, S. & Evans, M. G. 1946 *J. Poly. Sci.* **1**, 237.
 Dewar, M. J. S. 1949 *The electronic theory of organic chemistry*. Oxford: Clarendon Press.
 Eyring, H., Tobolsky, A. V., Harman, R. A. & Hulburt, H. M. 1943 *Ann. N.Y. Acad. Sci.* **44**, 371.

- Farmer, E. H. 1942 *Rubber Chem. & Technol.* **15**, 765.
 Flory, P. J. 1937 *J. Amer. Chem. Soc.* **59**, 241.
 Fieser, L. F. & Oxford, A. E. 1942 *J. Amer. Chem. Soc.* **64**, 2060.
 Glazebrook, H. H. & Pearson, T. G. 1936 *J. Chem. Soc.* p. 1777.
 Hey, D. H. & Waters, W. A. 1937 *Chem. Rev.* **21**, 169.
 Kharasch, M. S. & Gladstone, M. T. 1943 *J. Amer. Chem. Soc.* **65**, 15.
 Kharasch, M. S. & Brown, H. C. 1940 *J. Amer. Chem. Soc.* **62**, 925.
 Mayo, F. R. 1943 *J. Amer. Chem. Soc.* **65**, 2324.
 Mayo, F. R. 1948 *J. Amer. Chem. Soc.* **70**, 3689.
 Moore, J. W. & Taylor, H. S. 1940 *J. Chem. Phys.* **8**, 466.
 Price, C. C. 1946 *Reactions at carbon-carbon double bond*, p. 61. New York: Interscience Publishers, Inc.
 Price, C. C. & Schwarcz, M. 1940 *J. Amer. Chem. Soc.* **62**, 2891.
 Price, C. C. & Tate, B. E. 1943 *J. Amer. Chem. Soc.* **65**, 517.
 Remick, A. E. 1949 *Electronic interpretation of organic chemistry*, p. 137. New York: John Wiley and Sons, Inc.
 Schulz, G. V., Dinglinger, A. & Husemann, E. 1939 *Z. phys. Chem. B*, **43**, 385.
 Schulz, G. V. & Husemann, E. 1938 *Z. phys. Chem. B* **39**, 246.
 Suess, H., Pilch, K. & Rudorfer, H. 1937 *Z. phys. Chem. A*, **179**, 361.
 Suess, H. & Springer, A. 1937 *Z. phys. Chem. A*, **181**, 81.

Electronic levels in simple conjugated systems

I. Configuration interaction in cyclobutadiene

BY D. P. CRAIG

Sir William Ramsay and Ralph Forster Laboratories, University College, London

(Communicated by C. K. Ingold, F.R.S.—Received 9 March 1950)

In the simplest cyclic system of π -electrons, cyclobutadiene, a non-empirical calculation has been made of the effects of configuration interaction within a complete basis of antisymmetric molecular orbital configurations. The molecular orbitals are made up from atomic wave functions and all the interelectron repulsion integrals which arise are included, although those of them which are three- and four-centre integrals are only known approximately.

In this system configuration interaction is a large effect with a strongly differential action between states of different symmetry properties. Thus the $^1A_{1g}$ state is several electron-volts lower than the lowest configuration of that symmetry, whereas for $^1B_{1g}$ the comparable figure is about one-tenth of an electron-volt. The other two states examined, $^1B_{2g}$ and $^3A_{2g}$, are affected by intermediate amounts. The result is a drastic change in the energy-level scheme compared with that based on configuration wave functions.

Neither the valence-bond theory nor the molecular orbital theory (in which the four states have the same energy) gives a satisfactory account of the energy levels according to these results. One conclusion from the valence-bond theory which is, however, confirmed, is the somewhat unexpected one that the non-totally symmetrical $^1B_{2g}$ state is more stable than the totally symmetrical $^1A_{1g}$. On the other hand, it is clear that the valence-bond theory, with the usual value for its exchange integral, grossly exaggerates the resonance splitting of the states, giving separations between them several times too great. Thus the valence-bond theory leads to large values of the resonance energy (larger, per π -electron, than in benzene) and so associates with the molecule a considerable π -electron stabilization. This expectation has no support in the present more detailed and non-empirical calculations.

INTRODUCTION TO THREE PAPERS

This paper (I) and one soon to appear (II) by Coulson & Jacobs devoted to *trans*-butadiene, are in the first instance concerned with specific problems in their respective molecules. Their common principal concern, however, is to examine in simpler cases than hitherto the influence of configuration interaction in π -electron theory, and it appeared that a measure of common treatment would enlarge their usefulness and interest. This has led to the adoption, so far as possible, of a uniform notation and to the preparation of a joint summarizing discussion which appears as part III. This third paper draws together the more general issues raised by the two detailed ones and contains an assessment of the present position in the calculation of π -electron states as it now appears in this and some other recent work.

The following notation is used:

$I(\nu), II(\nu), \dots, K(\nu)$	Atomic $2p$ wave functions for electron ν centred on the 1st, 2nd, ..., k th nucleus.
W_{2p}	Energy of an atomic $2p$ electron.
$T(\nu)$	Kinetic energy operator for electron ν .
$H_K(\nu)$	The potential at ν due to a singly charged carbon ion, such that the relation $\{H_K(\nu) + T(\nu)\} K(\nu) = W_{2p} K(\nu)$ holds.
$\phi_r(\nu)$	r th normalized one-electron molecular orbital, numbered in a sequence of increasing energy.
ϵ_r	Energy of $\phi_r(\nu)$.
$\mathcal{H}_K(\nu)$	The potential at electron ν due to a neutral spherically symmetrical carbon atom.
H, E	The Hamiltonian (H) and its approximate eigenvalues (E), the latter suitably distinguished by subscripts or superscripts.

We use ψ and Ψ to denote configuration and state wave functions respectively. Subscripts and superscripts are attached to these in ways separately explained in the first two papers. Determinantal wave functions are abbreviated thus:

$$\sqrt{\frac{1}{4!}} \begin{vmatrix} \phi_l(1)\alpha(1) & \phi_m(1)\beta(1) & \phi_n(1)\alpha(1) & \phi_p(1)\beta(1) \\ \phi_l(2)\alpha(2) & & & \text{etc.} \\ \phi_l(3)\alpha(3) & & & \\ \phi_l(4)\alpha(4) & & & \end{vmatrix} \equiv (|\phi_l\bar{\phi}_m\phi_n\bar{\phi}_p|)$$

And, finally, the symmetry notation is that of Eyring, Walter, & Kimball (1944). Representations for one-electron wave functions are given small letters, and for molecular wave functions capitals.

The principal references are to Goeppert-Mayer & Sklar's (1938) original work with antisymmetric molecular orbital wave functions (ASMO), and to Parr & Crawford's (1948) recent compilation of the necessary energy integrals.

CYCLOBUTADIENE

Although the primary interest in cyclobutadiene in this paper is as a simple molecule in which to study the working of configuration interaction, it does, nevertheless, present in its own right an interesting problem to which non-empirical

calculations should be relevant. As Wheland (1938) first showed, the empirical† valence-bond (VB) and molecular orbital (MO) methods are in conflict in this molecule in a way they very rarely are. The VB method finds a *singlet* ground state and the MO finds a degenerate one, the *triplet* component of which comes to be lowest in the next higher approximation. This raises an important issue for the two simple methods, since the conviction they carry in calculations of π -electron ground states derives to a marked degree from the fact that, starting from strongly different assumptions, they are consistently in agreement, and any exception from this ought to be investigated.

The VB and MO approximations involve certain assumptions whose justification is not physical but simply that of analytical convenience. Thus it has been found convenient either to work with complete exclusion of polar structures (simple VB theory) or with an excessive inclusion of them (MO theory). Anything between these extremes needs a much more complicated formalism, as, for example, in adding some polar structures to the VB theory (Craig, 1950a). These features are of course always present, but their defects rarely emerge in energy calculations because the use of empirical constants tends to compensate for them in a consistent way. However, where the methods disagree, as they do in cyclobutadiene, an independent method of calculation is needed, and this will have to differ from the simple methods in two respects. It must avoid the special assumptions mentioned above by using more flexible wave functions, and it must use as little empirical material as possible. Both of these demands are met in the procedure used in these papers. It may be anticipated at the outset that such a method will be too complicated to have claims as a general method of approximation. Its value will be in dealing with particular questions in a non-empirical way, where empirical methods are for some reason unsatisfactory.

A suitable method has already been described in some detail (Craig 1950b). This has its starting point in the work of Goeppert-Mayer & Sklar (1938) (GMS) who gave a non-empirical means for getting the energies of antisymmetric molecular orbital wave functions in benzene. According to this, the electrons are supposed to move independently of one another in orbitals appropriate to the field of six carbon ions. The six electrons are assigned according to the Pauli principle to the lowest of these orbitals to give the correct symmetry for the state under consideration, and the energies of the occupied orbitals are added to give the 'orbital' or 'core' energy. Such an assignment of electrons to molecular orbitals, from its close analogy to the atomic case, is called a configuration. The total energy is the core energy together with the repulsion energy between the electrons moving in the orbitals to which they have been assigned. The GMS method thus combines the use of non-empirical energies with the use of determinantal wave functions of the MO type. It is possible, however, for the second feature, which imposes a highly special form on the wave

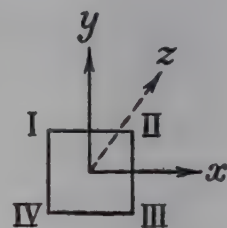
† The term empirical in this connexion has come to apply to methods which, dealing with a series of closely related molecules, give the results for all in terms of a common parameter. A value is assigned to the parameter by equating the calculated result for one member of the series to an experimental quantity for it; e.g. a thermochemical one. The Goeppert-Mayer & Sklar method, often called by contrast non-empirical, uses the experimental C-C distance but no other experimental molecular quantity.

functions, to be eliminated. Electron repulsion not only alters the energy of each configuration but also introduces a resonance between different configurations, so that a much more general wave function may be taken which is a linear combination of all configurations and the energy minimized by the variation method. Such approximate eigenfunctions are linear combinations of determinantal MO wave functions. But it is important to recognize that they do not retain the typical property of MO wave functions, which is that of splitting into factors each of which is the wave function for one electron. Thus an important feature of the procedure is that although MO wave functions are used as a basis their special features do not persist. Indeed, the linear combinations could, as a highly special case, be wave functions of the VB type, although it is not likely that those would prove to be the energy-minimized wave functions that the method is designed to select.

Orbitals and their symmetry properties

Cyclobutadiene is supposed to be a square molecule. Its one-electron orbitals are conveniently written

$$\phi_l(\nu) \equiv \phi_{-l}^*(\nu) = \frac{1}{(4\sigma_l)^{\frac{1}{2}}} \sum_K e^{2\pi i l k/4} K(\nu) \quad (l = 0, \pm 1, 2).$$



σ_l is a normalizing coefficient and l is a label which has the property that $|l|$ is the number of nodal planes perpendicular to the ring. Evidently ϕ_1 and ϕ_{-1} are degenerate, corresponding to alternative places for the single nodal plane. These orbitals and the wave functions for the molecule may be classified according to the representations of the group D_{4h} . For molecular wave functions only the representations symmetric to reflexion in the plane of the square occur, and for one-electron orbitals only those antisymmetric to that reflexion. In the following character table C_n means an n -fold rotation about the principal axis, C'_2 is a rotation about an in-plane axis passing through opposite atoms, and C''_2 is a rotation about an in-plane axis passing midway between opposite pairs of atoms.

CHARACTER TABLE FOR D_4

	E	C_2	$2C_4$	$2C'_2$	$2C''_2$	iE , for D_{4h}	
A_1	1	1	1	1	1	1	A_{1g}
						-1	A_{1u}
A_2	1	1	1	-1	-1	1	A_{2g}
						-1	A_{2u}
B_1	1	1	-1	1	-1	1	B_{1g}
						-1	B_{1u}
B_2	1	1	-1	-1	1	1	B_{2g}
						-1	B_{2u}
E	2	-2	0	0	0	2	E_g
						-2	E_u

The representations to which the one-electron orbitals belong are the following:

$$\phi_0 \subset a_{2u}, \quad \phi_{\pm 1} \subset e_g, \quad \phi_2 \subset b_{2u}.$$

Configuration wave functions and energies

Configuration wave functions belonging to any of the five plane-symmetric representations may now be set up. For example, the singlet A_{1g} wave function in the lowest configuration is, in an obvious notation,

$$\psi^I(A_{1g}) = 1/\sqrt{2}(|\phi_0\bar{\phi}_0\phi_1\bar{\phi}_{-1}| - |\phi_0\bar{\phi}_0\bar{\phi}_1\phi_{-1}|).$$

The Hamiltonian is that introduced by Goeppert-Mayer & Sklar (1938)

$$H = \sum_{\nu} T(\nu) + \sum_K \sum_{\nu} \left\{ \mathcal{H}_K(\nu) - \int \frac{e^2}{r_{\nu\sigma}} K^2(\sigma) d\tau_{\sigma} \right\} + \sum_{\alpha < \beta} \frac{e^2}{r_{\alpha\beta}}$$

in which the second expression in the bracket is the potential at electron ν due to an electron in a $2p$ orbital centred at nucleus K . The bracketed expression as a whole therefore is the attractive potential of a carbon ion, i.e. of a neutral atom minus its $2p_z$ electron, and is thus identical with the quantity $H_K(\nu)$ defined earlier. The symmetry of the system ensures that the ϕ_l are eigenfunctions, belonging to ϵ_l , of the Hamiltonian excluding electron repulsion, and this latter may be regarded as a perturbing potential. The energy of the A_{1g} configuration wave function is

$$\begin{aligned} E^I(A_{1g}) &= 2\epsilon_0 + 2\epsilon_1 + \iint \psi^{I*}(A_{1g}) \sum_{\alpha < \beta} \frac{e^2}{r_{\alpha\beta}} \psi^I(A_{1g}) d\tau_{\alpha} d\tau_{\beta} \\ &= 2\epsilon_0 + 2\epsilon_1 + \gamma_{00} + 4\gamma_{01} + \gamma_{11} - 2\delta_{01} + \delta_{1-1}, \end{aligned}$$

the first line having been expanded in the second into integrals over molecular orbitals. These are defined

$$\begin{aligned} \zeta(mn:pq) &= \iint \phi_m(\alpha) \phi_n(\alpha) \frac{e^2}{r_{\alpha\beta}} \phi_p(\beta) \phi_q(\beta) d\tau_{\alpha} d\tau_{\beta}, \\ \gamma_{ll'} &= \zeta(l, -l:l', -l'), \quad \delta_{ll'} = \zeta(l, -l':-l, l'). \end{aligned}$$

The coulomb (γ) and exchange (δ) integrals occur most commonly and their definitions agree with GMS. Complex conjugate wave functions have been cleared from the general form (ζ). The LCAO form of the molecular orbitals allows these integrals to be expressed as integrals over atomic $2p$ wave functions, the values of which are set out in the left-hand column of table 1. The notation is

$$I, II:III, IV = \iint I(\alpha) II(\alpha) \frac{e^2}{r_{\alpha\beta}} III(\beta) IV(\beta) d\tau_{\alpha} d\tau_{\beta}.$$

The values in table 1 apply to a C-C distance 1.40 Å. and to an effective nuclear charge $Z = 3.18$ electrons. The left-hand column excluding (I, I:III, III) contains values taken from Parr & Crawford (1948). (I, I:III, III) has been found from Griffing's (1947) interpolation formula for coulomb repulsions. It will be noticed that if we include all coulomb interactions, but exclude all but nearest neighbours otherwise, then the integrals already referred to, all of which come from integration of the appropriate expressions, are sufficient. If, on the other hand, we wish to include all the repulsions, then the remaining values in the left-hand column are needed. But, with two unimportant exceptions, these are three- and four-centre integrals, and their values

TABLE 1

Integrals over atomic orbitals (eV)		integrals over molecular orbitals (eV)		
			excluding three- and four-centre integrals	including all integrals
I, I:I, I	16.9310	$\epsilon_1 - \epsilon_0$	-6.3776	6.4233
I, I:II, II	9.0225	$\epsilon_2 - \epsilon_1$	1.1417	5.0027
I, I:III, III	6.7812	$\epsilon_2 - \epsilon_0$	-5.2359	11.4260
I, I:I, II	3.3136	γ_{00}	6.9214	10.7579
I, II:I, II	0.9522	γ_{01}	9.4027	10.3470
		γ_{02}	10.2531	10.3982
		γ_{11}	12.7033	10.2791
I, I:II, III	2.1510	γ_{12}	13.7067	10.4811
I, I:II, IV	1.0674	γ_{22}	14.4866	11.1288
I, I:I, III	1.0727	δ_{01}	4.3258	2.6467
I, II:I, III	0.3110	δ_{02}	1.5312	1.5369
I, II:I, IV	0.7719	δ_{1-1}	2.8827	1.7398
I, II:III, IV	0.6099	δ_{12}	-0.5771	2.9739
I, III:II, IV	0.1479	$\zeta(11:11)$	0.5654	1.6954
I, III:I, III	0.1479	$\zeta(01:12)$	2.3641	2.7137
		$\zeta(02:11)$	1.6248	1.6248

have to be obtained in a much less exact manner. The procedure in these cases (London 1945) may be illustrated by the following example:

$$I, I:II, III = \iint I^2(\alpha) \frac{e^2}{r_{\alpha\beta}} II(\beta) III(\beta) d\tau_\alpha d\tau_\beta.$$

The product $e II(\beta) III(\beta)$ is supposed equivalent to $SeK^2(\beta)$ (S is the overlap integral) for K placed midway between atoms II and III, and the interaction between this 'overlap' charge and the electron at atom I is worked out as a coulomb repulsion. Similar devices applied to the other integrals lead to the values given.

The energies of the configuration wave functions may now be written down. An exhaustive list of configurations and the types of wave function they admit is given in table 2. A configuration, for example, in which ϕ_0 is doubly occupied and ϕ_1 and ϕ_{-1} each singly occupied is abbreviated to $0^2, 1, -1$. The configurations in table 2 exhaust the possibilities for four electrons in the four orbitals and the set of wave functions is complete in this sense. There are 20 singlet wave functions and 15 triplet when degeneracy in the configuration $0, 1, -1, 2$ is allowed for.

TABLE 2

configurations	wave-function species
$0^2, 1^2; 0^2, -1^2, 0^2, 1, -1$	$^1B_{2g}; ^1B_{1g}; ^1A_{1g}; ^3A_{2g}$
$0^2, \pm 1, 2$	$1, ^3E_u$
$0, 1^2, -1; 0, -1^2, 1$	$1, ^3E_u$
$0, 1^2, 2; 0, -1^2, 2; 0, 1, -1, 2$	$1, ^3B_{1g}; 1, ^3A_{2g}; 1, ^3A_{1g}; 1, ^3, ^5B_{2g}$
$0^2, 2^2; -1^2, 1^2$	$^1A_{1g}; ^1A_{1g}$
$0, \pm 1, 2^2$	$1, ^3E_u$
$-1^2, 1, 2; 1^2, -1, 2$	$1, ^3E_u$
$1^2, 2^2; -1^2, 2^2; 1, -1, 2^2$	$^1B_{2g}; ^1B_{1g}; ^1A_{1g}; ^3A_{2g}$

The interaction of configurations

Attention is now confined to the species which arise in the lowest configuration, viz. $^1B_{2g}$, $^1A_{1g}$, $^1B_{1g}$ and $^3A_{2g}$. We shall give the calculation of the lowest A_{1g} state, allowing that to illustrate a procedure common to all. The calculation of configuration interaction is a variation calculation, using electron repulsion as perturbing potential over a variation function composed of all the configuration wave functions belonging to the same multiplicity and symmetry. In the approximation in which all repulsion integrals are included this leads to the following secular equation for the $^1A_{1g}$ states. Reading from left to right the columns refer to configurations in the order in which they appear in table 2. Because of the orthogonality of the one-electron molecular orbitals only the interelectron repulsion terms in the Hamiltonian contribute to the off-diagonal matrix components:

$$\begin{vmatrix} x & 0.7579 & 4.2057 & -3.7430 & 1.5369 \\ 0.7579 & x+12.4998 & -5.4274 & 5.4274 & 0.7579 \\ 4.2057 & -5.4274 & x+11.5397 & 0 & 3.7430 \\ -3.7430 & 5.4274 & 0 & x+12.1702 & -4.2057 \\ 1.5369 & 0.7579 & 3.7430 & -4.2057 & x+23.1049 \end{vmatrix} = 0$$

TABLE 3. SUMMARY OF RESULTS

excluding three- and four-centre integrals				including all integrals		
	energy† of lowest configuration wave-function	best value	depression	energy of lowest configuration wave-function	best value	depression
$^1A_{1g}$	51.47	46.79	4.67	58.87	55.57	3.30
$^1B_{2g}$	48.02	45.64	2.38	55.44	53.02	2.42
$^3A_{2g}$	45.70	44.64	1.06	55.39	53.73	1.66
$^1B_{1g}$	49.15	48.96	0.18	58.83	58.69	0.14

† In eV. To aid tabulation the common additive term $2\epsilon_0 + 2\epsilon_1$ is everywhere omitted.

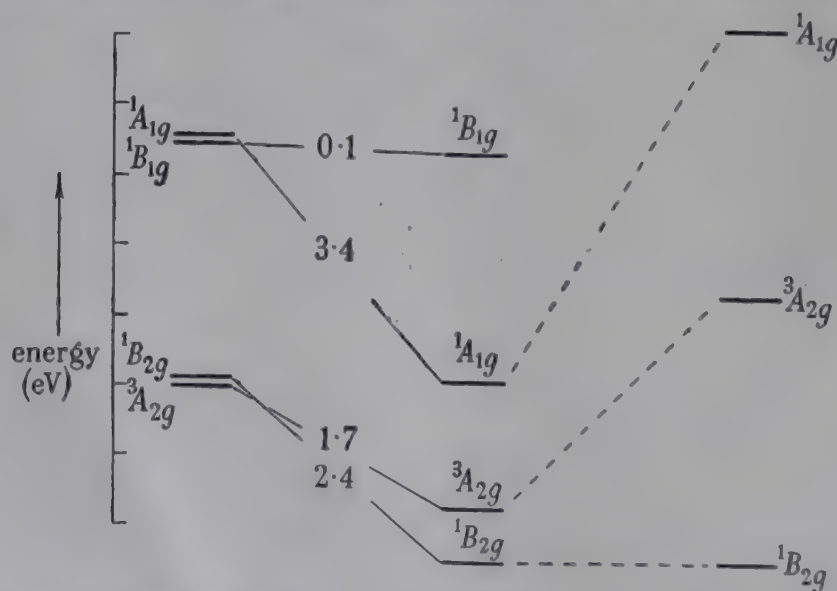


FIGURE 1. The states of cyclobutadiene (centre) are compared with the lowest configuration energies (left) and the VB values (right). Results are those from table 3 (right-hand side). The lowest VB state is located arbitrarily in the energy scale.

x stands for $2\epsilon_0 + 2\epsilon_1 + 58.8714 \text{ eV} - E$, where E is the energy variable. The lowest root leads to $E(^1A_{1g}) = 2\epsilon_0 + 2\epsilon_1 + 55.57 \text{ eV}$. This result and those for the other states, both with and without three- and four-centre integrals, are summarized in table 3.

DISCUSSION

Evidently in this molecule, as in benzene (Craig 1950*b*) and in naphthalene (Jacobs 1949), configuration interaction is a large effect, with a strong differential action between states of different symmetry properties. Thus there is almost no change in the $^1B_{1g}$ state, while the $^1A_{1g}$ is depressed by several electron-volts. This has the effect of reversing the order of these two states. It may be expressed otherwise by saying that the lowest configuration wave function for $^1B_{1g}$ is found to be very nearly the best wave function that can be formed from $2p$ atomic wave functions, whereas for $^1A_{1g}$ the lowest configuration wave function is a very poor approximation needing considerable admixture of the configurations denoted earlier as 0^2 , 2^2 and -1^2 , 1^2 to minimize the energy. The $^1B_{2g}$ and $^3A_{2g}$ states fall between these extremes, being depressed by 1 to 2 eV. It may in any case be argued generally that a triplet configuration wave function is less prone to improvement by configuration interaction than the associated singlet (in this case $^1A_{1g}$ in table 2) because a degree of correlation between electrons is already achieved in it by the fact that the two unpaired electrons cannot occupy the same atomic orbital.

The calculations for this simple molecule have been made to include all configurations, but this will be possible in few, if any, larger molecules. It would be useful, therefore, to have the means for selecting beforehand the configurations whose interactions play the dominating part in depressing the energy. This is a considerable difficulty. Obviously only configurations having matrix elements with the lowest configuration need be considered in the first instance, but to make a selection from these is shown to be hazardous by the fact that the off-diagonal matrix components of the energy are very sensitive to the inclusion of (and therefore to the precision of) many-centre integrals. Thus it is not likely that any general rule can be found for selecting the important configurations. This point is nicely illustrated in the $^1A_{1g}$ calculations. When many-centre integrals are omitted the configuration 0^2 , 2^2 interacts very little with the lowest configuration, there being a matrix element of only 0.8 eV between the two. When many-centre integrals are included the matrix component is 4.2 eV and the configuration 0^2 , 2^2 plays an important part in depressing the energy. Where, as here, the calculation extends over all configurations, these variations in particular matrix components do not imply large variations in the final depressions (see table 3), but this will not be true when it is necessary to select out of a large number of configurations a few upon which to base a computation.

In comparing the results with those from the empirical theories two conclusions may be drawn without presuming higher accuracy than the non-empirical work is likely to give. It will be remembered that according to the VB theory the $^1B_{2g}$ state is lowest, and the $^3A_{2g}$ and $^1A_{1g}$ states fall above it in that order in steps of 2α , α being the exchange integral with the value 1.92 eV in benzene. According to the MO theory all four states we have considered fall together at the same energy.

Evidently neither of these pictures is satisfactory according to the present work. But in the first place it is confirmed that the non-totally symmetrical state $^1B_{2g}$ is more stable than the totally symmetrical $^1A_{1g}$. This is somewhat unexpected, since in all the known conjugated hydrocarbons the lowest singlet state belongs to the totally symmetrical representation, and it might have been expected that this would be true of cyclobutadiene. In a many-electron system, however, there appears to be no theoretical support for this expectation.

The second and more important point is that the VB method, with parameter chosen to fit the simple known hydrocarbons, appears to overestimate grossly the intervals between the $^1B_{2g}$, $^3A_{2g}$ and $^1A_{1g}$ states of cyclobutadiene. The separation between the first two according to the present calculations using all interaction integrals is 0.7 eV as compared with the VB value 3.8 eV. The interval between the first and third, which in the VB theory gives directly the measure of the 'resonance' between the two Kekulé structures for cyclobutadiene, is 2.5 eV compared with 7.6 eV in the VB theory itself. The precision of the calculations does not allow a final choice between $^1B_{2g}$ and $^3A_{2g}$ as the lowest state, but it is clear that these two in any case fall closely together, and that the scale of resonance interactions is very much smaller than suggested by the VB theory with its usual parameter; that is to say, the resonance interaction is much smaller than it is in the known hydrocarbons from which this parameter is worked out. It follows, of course, that the VB theory also greatly exaggerates the *resonance energy* of cyclobutadiene, i.e. the amount by which the molecule is expected to be more stable than one of its Kekulé structures.

This work has been carried out during the tenure of a Turner and Newall Research Fellowship in the University of London.

REFERENCES

- Craig, D. P. 1950a *Proc. Roy. Soc. A*, **200**, 401.
Craig, D. P. 1950b *Proc. Roy. Soc. A*, **200**, 474.
Eyring, H., Walter, J. & Kimball, E. 1944 *Quantum chemistry*. New York: John Wiley.
Goeppert-Mayer, M. & Sklar, A. L. 1938 *J. Chem. Phys.* **6**, 645.
Griffing, V. 1947 *J. Chem. Phys.* **15**, 421.
Jacobs, Miss J. 1949 *Proc. Phys. Soc.* **42**, 710.
London, A. 1945 *J. Chem. Phys.* **13**, 306.
Parr, R. G. & Crawford, B. 1948 *J. Chem. Phys.* **16**, 1049.
Wheland, G. W. 1938 *Proc. Roy. Soc. A*, **164**, 397.

Asymptotic expansions for the hypergeometric functions occurring in gas-flow theory

BY T. M. CHERRY

(Communicated by S. Goldstein, F.R.S.—Received 10 March 1950)

Asymptotic formulae of two sorts are found for the hypergeometric functions, of large order ν , which arise when gas-flows are investigated by the hodograph method. These functions are monotonic for 'subsonic' values of the argument τ , oscillating for 'supersonic' values. (i) Elementary formulae, involving exponential or circular functions and singular at the sonic value $\tau = \tau_s$, are developed to the terms in ν^{-4} . When τ is near τ_s these formulae give poor approximation. (ii) Formulae giving uniform approximation near τ_s , involving Bessel functions, are developed to the terms in ν^{-4} . These give 7-figure accuracy over the physically significant range of τ , for $|\nu| \geq 10$. The coefficients are tabulated for $\tau = 0$ (0.02) 0.50 for the case where the adiabatic index γ of the gas is 1.4.

1. INTRODUCTION. DEFINITIONS

In the hodograph theory for investigating gas-flows in two dimensions, solutions for the Legendre potential are obtained in the form

$$\Sigma A_\nu y(\nu, \tau) e^{i\nu\theta}, \quad (1.1)$$

where $y(\nu, \tau)$ is a solution of the differential equation

$$(\chi)^\dagger: \quad \frac{d^2 y}{d\tau^2} + \left(\frac{1}{\tau} - \frac{\beta}{1-\tau} \right) \frac{dy}{d\tau} + \nu^2 \left(\frac{\beta^2}{2\tau(1-\tau)} - \frac{1}{4\tau^2} \right) y = 0. \quad (1.2)$$

There are also series of the form (1.1) for the stream function, the differential equation for y being now

$$(\psi)^\dagger: \quad \frac{d^2 y}{d\tau^2} + \left(\frac{1}{\tau} + \frac{\beta}{1-\tau} \right) \frac{dy}{d\tau} + \nu^2 \left(\frac{\beta^2}{2\tau(1-\tau)} - \frac{1}{4\tau^2} \right) y = 0. \quad (1.3)$$

The physical meaning of the symbols is that $\tau^{\frac{1}{2}} \cos \theta$, $\tau^{\frac{1}{2}} \sin \theta$ are rectangular velocity components, and $\beta = 1/(\gamma - 1)$, where γ is the adiabatic index of the gas.

The series (1.1) does not usually converge rapidly, and it is therefore important that we know the analytical and numerical behaviour of the solutions of the differential equations (1.2), (1.3) when ν is large. Such knowledge is to be obtained from their asymptotic expansions, and it is with these that this paper is concerned. The properties of the two differential equations are so similar that it will be sufficient to give the exposition for (1.2) and merely state the results for (1.3). We shall restrict τ to the physically significant range $0 \leq \tau < 1$, but shall allow ν to be complex. The theory is valid for any positive value of the constant β .

The coefficient of y in (1.2) is negative for $0 < \tau < \tau_s$, positive for $\tau_s < \tau < 1$, where

$$\tau_s = 1/(1 + 2\beta). \quad (1.4)$$

(The gas-flow is subsonic for $0 < \tau < \tau_s$, supersonic for $\tau_s < \tau < 1$.) Hence the solutions are monotonic in the first range, but oscillating in the second; and, as concerns

† The symbols χ , ψ are mnemonics; see (1.7), 1.11).

approximations in terms of elementary functions (analogous to Debye's 'A' series for Bessel functions of large order),[†] separate asymptotic expansions are required in the two ranges. Their coefficients will be calculated up to the terms in $1/\nu^4$ (§ 2). These 'elementary' expansions give non-uniform approximation near the transition point τ_s , and there is in fact a wide range of τ in which (numerically speaking) the expansions are useless unless ν is prohibitively large. In § 3 are given asymptotic expansions of a new type which are free from this defect. The coefficients are calculated up to the terms in $1/\nu^4$, and are tabulated (table 2) for $\gamma = 1.4$. From this table and the appropriate formulae of § 3 the solutions of (1.2), (1.3) can be calculated correct to 7 significant figures, when $|\nu| \geq 10$.

For the sort of asymptotic expansion dealt with in § 2 the general lines of the theory are well known; and the leading approximations to the primary solutions $\chi_\nu(\tau)$, $\psi_\nu(\tau)$ have been given by a number of authors (Tsien & Kuo 1946; Seifert 1947; Lighthill 1947; Cherry 1947). Hence the weight of § 2 lies in the discussion of the full asymptotic series and of the logarithmic solutions.

Definition of the solutions of (1.2). THEOREM. Let

$$2a_\nu = \nu - \beta + \sqrt{\{\nu^2(1 + 2\beta) + \beta^2\}}, \quad 2b_\nu = \nu - \beta - \sqrt{\{\nu^2(1 + 2\beta) + \beta^2\}} \quad (1.5)$$

and
$$h_\nu = \frac{\Gamma(1 + a_\nu) \Gamma(\nu - b_\nu)}{\Gamma(1 + a_\nu - \nu) \Gamma(-b_\nu) \Gamma(1 + \nu) \Gamma(\nu)}. \quad (1.6)$$

Then, if ν is not a negative integer, (1.2) has the solutions

$$\chi_\nu(\tau) = \tau^{1/2} F(\nu - a_\nu, \nu - b_\nu; \nu + 1; \tau), \quad (1.7)$$

$$\chi y_\nu(\tau) = \chi_\nu(\tau) \cot \nu\pi - \chi_{-\nu}(\tau)/(\pi h_\nu), \quad (1.8)$$

where F denotes the hypergeometric series; for n positive integral, $\chi y_n(\tau)$ is defined by the limiting value of (1.8) for $\nu \rightarrow n$, which exists.[‡]

If ν is not integral there is also the solution $\chi_{-\nu}(\tau)$.

Note 1. By (1.5), a_ν, b_ν are the two determinations of a double-valued function, branched at $\nu = \pm i\beta(1 + 2\beta)^{-1/2}$; these are distinguished by the convention that the ν -plane is cut between these branch points and that for $|\arg \nu| < \frac{1}{2}\pi$ the square root has its principal value; so that ν real positive makes a_ν positive and b_ν negative. The distinction between a_ν, b_ν is relevant for the definitions of $h_\nu, \chi y_\nu$; but χ_ν involves a_ν, b_ν symmetrically, and is single-valued in the uncut ν -plane; it has simple poles at $\nu = -1, -2, \dots$

Note 2. The solutions of (1.2) have analogies with Bessel functions, which will appear most clearly in the formulae of § 3. For ν non-integral there are solutions $\chi_\nu, \chi_{-\nu}$ analogous to $J_\nu, J_{-\nu}/\sin \nu\pi$. When ν is the positive integer n , χ_{-n} is infinite and a new, logarithmic, solution is required. The 'best' choice of this depends upon the context; as regards asymptotic formulae it is χy_n as defined (and named from

[†] See, for example, Watson's *Bessel Functions*, §§ 8.4 to 8.42.

[‡] The existence of this limit will be shown in § 2, Theorem 2; the other assertions are trivial.

analogy with the Bessel function Y_n in (1.8). In other contexts the following combinations of $\chi_n, \chi y_n$ are important:

$$\chi_{-n}^*(\tau) = \lim_{\nu \rightarrow -n} \left\{ \chi_\nu(\tau) + \frac{h_n \chi_{-\nu}(\tau)}{\nu + n} \right\} \quad (1.9)$$

$$= -\pi h_n \chi y_n(\tau) + h'_n \chi_n(\tau), \quad h'_n = [dh_\nu/d\nu]_{\nu=n}; \quad (1.9)^*$$

and

$$\chi_{-n}^{**}(\tau) = -\pi h_n \chi y_n(\tau) + h_n \{ \Psi(1 + a_n - n) + \Psi(n - b_n) - \Psi(n + 1) - \Psi(1) \} \chi_n(\tau) \quad (1.10)$$

$$\begin{aligned} &= \tau^{-\frac{1}{2}n} \left[1 - \frac{a_n b_n}{n-1} \frac{\tau}{1!} + \frac{a_n(a_n-1)b_n(b_n-1)}{(n-1)(n-2)} \frac{\tau^2}{2!} - \dots \right. \\ &\quad \left. + \frac{a_n \dots (a_n - n + 2) b_n \dots (b_n - n + 2) (-\tau)^{n-1}}{(n-1) \dots 1 (n-1)!} \right] - h_n \chi_n(\tau) \log \tau \\ &\quad + \sum_{r=1}^{\infty} \frac{\Gamma(1+a_n) \Gamma(n-b_n+r)}{\Gamma(-b_n) \Gamma(1+a_n-n-r)} \frac{(-1)^r \tau^{\frac{1}{2}n+r}}{(n-1)! (n+r)! r!} [\Psi(1+a_n-n) + \Psi(n-b_n) \\ &\quad - \Psi(n+1) - \Psi(1) - \Psi(1+a_n-n-r) - \Psi(n-b_n+r) \\ &\quad + \Psi(n+1+r) + \Psi(1+r)], \quad (1.10)^* \end{aligned}$$

where $\Psi(z) = \Gamma'(z)/\Gamma(z)$. For $\tau \sim 0$ both of these have the form $\tau^{-\frac{1}{2}n}\{1 + O(\tau)\}$, and the notation $\chi_{-n}^*, \chi_{-n}^{**}$ has been chosen from this resemblance to the general χ_ν . The significance of χ_{-n}^{**} is that it is the only combination of χy_n and χ_n whose ascending series contains no term in $\tau^{\frac{1}{2}n}$ and which $\sim \tau^{-\frac{1}{2}n}$ for $\tau \sim 0$.

Definition of solutions of (1.3). Let a_ν, b_ν, h_ν be defined by (1.5), (1.6). Then (1.3) has the solutions

$$\psi_\nu(\tau) = \tau^{\frac{1}{2}\nu} F(a_\nu, b_\nu; \nu + 1; \tau), \quad (1.11)$$

$$\psi y_\nu(\tau) = \psi_\nu(\tau) \cot \nu\pi - \frac{(\nu+1) \psi_{-\nu}(\tau)}{\pi(\nu-1) h_\nu}. \quad (1.12)$$

For $n = 2, 3, \dots$, but not for $n = 1$, $\psi y_\nu(\tau)$ has a limit $\psi y_n(\tau)$ as $\nu \rightarrow n$; this is a logarithmic solution, and other such solutions are:

$$\psi_{-n}^*(\tau) = \lim_{\nu \rightarrow -n} \left\{ \psi_\nu(\tau) + \frac{(n-1) h_n \psi_{-\nu}(\tau)}{(n+1)(\nu+n)} \right\} \quad (1.13)$$

$$= -\frac{\pi(n-1) h_n}{n+1} \psi y_n(\tau) + \left\{ \frac{2h_n}{(n+1)^2} + \frac{(n-1) h'_n}{n+1} \right\} \psi_n(\tau), \quad (1.13)^*$$

$$\psi_{-n}^{**}(\tau) = -\frac{\pi(n-1) h_n}{n+1} \psi y_n(\tau) + \frac{(n-1) h_n}{n+1} \{ \Psi(a_n) + \Psi(1-b_n) - \Psi(n+1) - \Psi(1) \} \psi_n(\tau) \quad (1.14)$$

$$\begin{aligned} &= \tau^{-\frac{1}{2}n} \left[1 - \frac{(a_n-n)(b_n-n)}{n-1} \frac{\tau}{1!} + \frac{(a_n-n)(a_n-n+1)(b_n-n)(b_n-n+1)}{(n-1)(n-2)} \frac{\tau^2}{2!} - \dots \right. \\ &\quad \left. + \frac{(a_n-n) \dots (a_n-2)(b_n-n) \dots (b_n-2) (-\tau)^{n-1}}{(n-1) \dots 1 (n-1)!} \right] - \frac{(n-1) h_n}{n+1} \psi_n(\tau) \log \tau \\ &\quad + \sum_{r=1}^{\infty} \frac{\Gamma(a_n+r) \Gamma(n+1-b_n)}{\Gamma(a_n-n) \Gamma(1-b_n-r)} \frac{(-1)^r \tau^{\frac{1}{2}n+r}}{(n-1)! (n+r)! r!} [\Psi(a_n) + \Psi(1-b_n) - \Psi(n+1) \\ &\quad - \Psi(1) - \Psi(a_n+r) - \Psi(1-b_n-r) + \Psi(n+1+r) + \Psi(1+r)]. \quad (1.14)^* \end{aligned}$$

[It will be noted that $(\nu - 1)h_\nu/(\nu + 1)$ plays the same part in ψ -formulae as h_ν does in χ -formulae. Hence arises a difference regarding $\nu = 1$; here χ_{-1} is singular and χy_ν regular, whereas ψ_{-1} is regular and ψy_ν singular. It is easily proved that

$$\psi_{-1}(\tau) = \tau^{-\frac{1}{2}} + \frac{1}{2}\beta\psi_1(\tau).]$$

Other notations. Investigations concerning one or both of the differential equations (1.2), (1.3) have been published by perhaps a dozen different authors, and there are almost as many different notations for their solutions. Most authors have worked with symbols for hypergeometric functions rather than functions like $\chi_\nu(\tau)$ in which the factor $\tau^{\frac{1}{2}\nu}$ is incorporated. But by using symbols like χ_ν we simplify the formulae relating to gas-flows and also those relating to asymptotic expansions.

It seems unnecessary to list the notations of other authors for the χ_ν , ψ_ν or $\tau^{-\frac{1}{2}\nu}\chi_\nu$, $\tau^{-\frac{1}{2}\nu}\psi_\nu$ of this paper; when regard is had to the particular context there will be no doubt as to the identification. But there is no unique choice of logarithmic solutions, and the following information may be useful:

Garrick & Kaplan (1944) use the symbol $Y_\nu(\tau)$ with two different meanings. For values of ν other than $-2, -3, \dots$, $Y_\nu(\tau) = \tau^{-\frac{1}{2}\nu}\psi_\nu(\tau)$. For $n = 2, 3, \dots$,

$$Y_{-n}(\tau) = \tau^{\frac{1}{2}n}\psi_{-n}^{**}(\tau).$$

Huckel (1948) gives tables of Garrick & Kaplan's Y_ν .

Tsien & Kuo (1946) define $\tilde{G}_\nu$ and G_ν :

$$\tilde{G}_\nu(\tau) = \tau^{\frac{1}{2}\nu}h_\nu\{-\pi\chi y_\nu(\tau) + \cos\pi(\nu - a_\nu)\operatorname{cosec}^2\pi(\nu - a_\nu)\chi_\nu(\tau)\},$$

$$G_\nu(\tau) = \frac{\tau^{\frac{1}{2}\nu}(\nu - 1)h_\nu}{\nu + 1}\{-\pi\psi y_\nu(\tau) + \cos\pi b_\nu\operatorname{cosec}^2\pi b_\nu\psi_\nu(\tau)\}.$$

These functions are not well chosen, since when ν is positive integral $\nu - a_\nu, b_\nu$ can be arbitrarily near to negative integers and the functions are then unbounded.

Lighthill (1947) defines ψ_n^* :

$$\psi_n^*(\tau) = \psi_{-n}^*(\tau) - \frac{(n-1)h_n}{n(n+1)}\psi_n(\tau).$$

Cherry (1949) defines

$$F_{-n}(\tau) = \tau^{\frac{1}{2}n}\chi_{-n}^*(\tau), \quad \tilde{F}_{-n}(\tau) = \tau^{\frac{1}{2}n}\psi_{-n}^*(\tau).$$

In these identifications, the symbols of this paper are used on the right.

Identities. If any two of $\chi_\nu, \chi'_\nu, \psi_\nu, \psi'_\nu$ are known, the others can be found from

$$(\nu - 1)(1 - \tau)^{-\beta}\psi_\nu(\tau) = \nu\chi_\nu(\tau) - 2\tau\chi'_\nu(\tau)/\nu, \quad (1.15)$$

$$(\nu + 1)(1 - \tau)^\beta\chi_\nu(\tau) = \nu\psi_\nu(\tau) + 2\tau\psi'_\nu(\tau)/\nu. \quad (1.16)$$

2. ELEMENTARY ASYMPTOTIC FORMULAE

THEOREM 1. *Let*

$$\alpha = \frac{1}{2}(1 + 2\beta)^{\frac{1}{2}} - \frac{1}{2} = \frac{1}{2}\tau_s^{-\frac{1}{2}} - \frac{1}{2}, \quad \delta = \alpha^\alpha(1 + \alpha)^{-1-\alpha}. \quad (2.1)$$

Then there are expansions in descending powers of ν :

$$\left. \begin{aligned} a_\nu &= \nu(1 + \alpha) - \frac{1}{2}\beta + \frac{1}{4}\beta^2/\nu(1 + 2\alpha) + \dots, \\ b_\nu &= -\nu\alpha - \frac{1}{2}\beta - \frac{1}{4}\beta^2/\nu(1 + 2\alpha) + \dots, \end{aligned} \right\} \quad (2.2)$$

and

$$\log(2\pi h_\nu) \sim -2\nu \log \delta + c_1/\nu + c_3/\nu^3 + \dots, \quad (2.3)$$

where

$$\begin{aligned} c_1 &= \frac{(1-\tau_s)^2}{4\tau_s^{\frac{3}{2}}} \operatorname{arc} \tanh \tau_s^{\frac{1}{2}} - \frac{(1-3\tau_s)(3-7\tau_s)}{12\tau_s(1-\tau_s)}, \\ c_3 &= -\frac{(1-\tau_s)^4}{64\tau_s^{\frac{3}{2}}} \operatorname{arc} \tanh \tau_s^{\frac{1}{2}} + \frac{(1-\tau_s)^4}{64\tau_s^2} + \frac{(1-\tau_s)^3}{192\tau_s} \\ &\quad - \frac{1}{32(1-\tau_s)^3} \left(\frac{67}{45} + \frac{223}{15}\tau_s - \frac{202}{5}\tau_s^2 + 10\tau_s^3 + 3\tau_s^4 - \frac{1}{3}\tau_s^5 \right); \end{aligned} \quad (2.4)$$

(2.2) are convergent for $|\nu| > \beta/(1+2\alpha)$; (2.3) is asymptotic for $\nu \sim \infty$ provided $|\arg \nu| < \pi$;† and the powers of ν are all odd except for the terms $-\frac{1}{2}\beta$ in (2.2).

Proof. (2.2) are elementary expansions of (1.5). To obtain (2.3) we expand the gamma functions in (1.6) by Stirling's theorem and then, inserting the expansions of $\nu - b_\nu$, etc., arrange the resulting double series in powers of ν . The absence of a term in $\nu \log \nu$ is directly seen, and the absence of terms in $\nu^0, \nu^{-2}, \dots$ is proved as follows: From (2.2) we have

$$a_{-\nu} = b_\nu - \nu, \quad b_{-\nu} = a_\nu - \nu, \quad (2.5)$$

and thence (1.6) gives

$$4\pi^2 h_\nu h_{-\nu} = \cot a_\nu \pi \cot a_{-\nu} \pi \left(1 - \frac{\cos(a_\nu - 2\nu)\pi}{\cos a_\nu \pi} \right) \left(1 - \frac{\cos(a_{-\nu} + 2\nu)\pi}{\cos a_{-\nu} \pi} \right).$$

Here

$$\begin{aligned} a_\nu &= \nu(1+\alpha) + \dots, & a_{-\nu} &= -\nu(1+\alpha) + \dots, & a_\nu - 2\nu &= -\nu(1-\alpha) + \dots, \\ a_{-\nu} + 2\nu &= \nu(1-\alpha) + \dots, \end{aligned}$$

and since β is positive, (2.1) gives α positive. Hence for $\operatorname{Im} \nu$ large positive,

$$4\pi^2 h_\nu h_{-\nu} = 1 + O(e^{2i\alpha\nu}) \quad \text{or} \quad 1 + O(e^{2i\nu}),$$

according as $\alpha < 1$ or $\alpha > 1$; and $\log(4\pi^2 h_\nu h_{-\nu})$ is exponentially small. The expansion of the form $\log(2\pi h_\nu) \sim \sum c_r \nu^{-r}$ for $|\arg \nu| < \pi$ implies $\log(4\pi^2 h_\nu h_{-\nu}) \sim 2(c_0 + c_2 \nu^{-2} + \dots)$ for $0 < \arg \nu < \pi$, and the preceding gives $c_0 = c_2 = \dots = 0$.

THEOREM 2. When $\operatorname{Re} a_\nu > -1$ and $\operatorname{Re} b_\nu < 0$,

$$\begin{aligned} \chi_\nu(\tau) &= \int_1^{(0-)} \frac{\tau^{1\nu} \Gamma(1+a_\nu-\nu) \Gamma(1+\nu)}{2i\pi \Gamma(1+a_\nu)} (-t)^{\nu-a_\nu-1} (1-t)^{a_\nu} (1-\tau t)^{b_\nu-\nu} dt \\ &= \int_1^{(0-)} \phi(t, \tau, \nu) dt, \quad \text{say}; \end{aligned} \quad (2.6)$$

$$\left. \begin{aligned} \frac{\chi_{-\nu}(\tau)}{2i\pi h_\nu} - \frac{e^{-\nu\pi i} \chi_\nu(\tau)}{2i \sin \nu\pi} &= \int_{i\infty}^1 \phi(t, \tau, \nu) dt, \\ -\frac{\chi_{-\nu}(\tau)}{2i\pi h_\nu} + \frac{e^{\nu\pi i} \chi_\nu(\tau)}{2i \sin \nu\pi} &= \int_1^{-i\infty} \phi(t, \tau, \nu) dt, \end{aligned} \right\} \quad (2.7)$$

where $(-t)^{\nu-a_\nu-1}$, etc., denote principal values.

† Of course, the range of $\arg \nu$ for uniformity of (2.3) is $|\arg \nu| \leq \pi - \epsilon$, for any positive ϵ . The ranges for (2.9), etc. are similarly given in the non-uniform form.

Note. By (2.2) the conditions $\text{Re } a_\nu > -1$, $\text{Re } b_\nu < 0$ are satisfied when

$$|\arg \nu| \leq \frac{1}{2}\pi - \epsilon$$

and $|\nu|$ is sufficiently large. But by trivial modifications we can convert (2.6), (2.7) into forms valid for $\arg \nu$ between $\frac{1}{2}\pi - \epsilon$, $\pi - \epsilon$ or between $-\frac{1}{2}\pi + \epsilon$, $-\pi + \epsilon$. The modification consists in using paths of integration that approach the limits 1 and ∞ on suitably chosen spirals. For example, if $\arg \nu$ is between $\frac{1}{2}\pi$, π the point 1 is approached on a spiral $|1-t| \exp \{c \arg (1-t)\} = \text{const.}$, where $0 < c < \text{Im } \nu / (-\text{Re } \nu)$.

Proof. By substitution in the differential equation (1.2) it is verified that the integrals on the right of (2.6), (2.7) are solutions, so if ν is not integral each has the value $A\chi_\nu(\tau) + B\chi_{-\nu}(\tau)$, where A, B are certain constants. These constants are found by taking limits for $\tau \rightarrow 0, 1$; the integrals become beta-function integrals, while for $\chi_\nu(1)$, $\chi_{-\nu}(1)$ we use the well-known formula for $F(A, B; C; 1)$. To use this formula we require $\text{Re}(a_\nu + b_\nu + 1 - \nu) > 0$, i.e. by (1.5) $\beta < 1$; but as a last step this restriction is removed by analytic continuation.

Finally, the integrals in (2.6), (2.7) are regular for ν positive integral, so here the expressions on the left of (2.7) must have limits and all the formulae remain valid. Also the difference of the two functions in (2.7) is $i\chi y_\nu(\tau)$, so the existence of this is proved for ν positive integral.

From analogy with the Hankel functions we introduce the notation

$$\left. \begin{aligned} \chi h_\nu^{(1)}(\tau) &= \frac{\chi_{-\nu}(\tau)}{i\pi h_\nu} - \frac{e^{-\nu\pi i}\chi_\nu(\tau)}{i \sin \nu\pi} = \chi_\nu(\tau) + i\chi y_\nu(\tau), \\ \chi h_\nu^{(2)}(\tau) &= -\frac{\chi_{-\nu}(\tau)}{i\pi h_\nu} + \frac{e^{\nu\pi i}\chi_\nu(\tau)}{i \sin \nu\pi} = \chi_\nu(\tau) - i\chi y_\nu(\tau). \end{aligned} \right\} \quad (2.8)$$

It is of interest, but not relevant for what follows, that $\chi h_\nu^{(1)}(\tau)$, $\chi h_\nu^{(2)}(\tau)$ are two different continuations into $|\tau| < 1$ of

$$C_\nu \tau^{b_\nu - \frac{1}{2}\nu} F(\nu - b_\nu, -b_\nu; 1 + a_\nu - b_\nu; \tau^{-1}),$$

where C_ν depends on ν only.

THEOREM 3. For $0 < \tau < \tau_s$, $\chi_\nu(\tau)$ has an asymptotic expansion

$$\chi_\nu(\tau) \sim \frac{(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}}{(1-\tau/\tau_s)^{\frac{1}{2}}} \delta^\nu e^{-\nu u} \left(1 + \frac{p_1}{\nu} + \frac{p_2}{\nu^2} + \dots \right) \quad (2.9)$$

valid for $\nu \sim \infty$ with $|\arg \nu| < \pi$, where δ is defined in (2.1),

$$u = \text{arc tanh} \sqrt{\left(\frac{1-\tau/\tau_s}{1-\tau} \right)} - \tau_s^{-\frac{1}{2}} \text{arc tanh} \sqrt{\frac{\tau_s - \tau}{1-\tau}}, \quad (2.10)$$

and $p_1, p_2, \dots$ are functions of τ defined as follows: Put

$$\zeta = \sqrt{\frac{1-\tau/\tau_s}{1-\tau}} \quad (2.11)$$

$$\text{and} \quad \phi = \frac{1-\zeta^2}{4(1-\tau_s)^2} \left\{ \frac{5}{\zeta^6} - \frac{1+6\tau_s}{\zeta^4} + \frac{3-4\tau_s}{\zeta^2} + (1-2\tau_s)(1-4\tau_s) \right\}. \quad (2.12)$$

Then

$$\left. \begin{aligned} p_1 &= \frac{1}{2} \int_0^\tau \phi \frac{du}{d\tau} d\tau = \frac{(1-\tau_s)^2}{8\tau_s^{\frac{3}{2}}} \{ \operatorname{arc} \tanh \zeta \tau_s^{\frac{1}{2}} - \operatorname{arc} \tanh \tau_s^{\frac{1}{2}} \} \\ &\quad + \frac{1}{8(1-\tau_s)} \left\{ -\frac{5}{3\zeta^3} + \frac{1+\tau_s}{\zeta} - \frac{(1-2\tau_s)(1-4\tau_s)}{\tau_s} \zeta + \frac{(1-3\tau_s)(3-7\tau_s)}{3\tau_s} \right\}, \\ p_2 &= \frac{1}{2} p_1^2 + \frac{1}{4} \phi, \\ p_3 &= \frac{1}{4} \phi p_1 + \frac{1}{8} \dot{\phi} + \frac{1}{6} p_1^3 + \frac{1}{8} \int_0^\tau \phi^2 \frac{du}{d\tau} d\tau, \\ p_4 &= \frac{1}{8} \phi^2 + \frac{1}{16} \ddot{\phi} + p_1 p_3 - \frac{1}{4} \phi p_1^2 - \frac{1}{4} p_1^4, \end{aligned} \right\} \quad (2.13)$$

where

$$\dot{\phi} = \frac{d\phi}{du} = \frac{(1-\zeta^2)(1-\tau_s\zeta^2)}{(1-\tau_s)\zeta^2} \frac{d\phi}{d\zeta}, \quad \ddot{\phi} = \frac{d^2\phi}{du^2},$$

$$\begin{aligned} \frac{1}{8} \int_0^\tau \phi^2 \frac{du}{d\tau} d\tau &= \frac{1}{128(1-\tau_s)^3} \left\{ -\frac{25}{9\zeta^9} + \frac{5(1+\tau_s)}{\zeta^7} - \frac{41-3\tau_s+\tau_s^2}{5\zeta^5} + \frac{27+19\tau_s-89\tau_s^2-\tau_s^3}{3\zeta^3} \right. \\ &\quad - \frac{3+37\tau_s-75\tau_s^2-7\tau_s^3+\tau_s^4}{\zeta} + \frac{1-7\tau_s+20\tau_s^2-52\tau_s^3+96\tau_s^4-64\tau_s^5}{\tau_s^2} \zeta \\ &\quad - \frac{1-12\tau_s+52\tau_s^2-96\tau_s^3+64\tau_s^4}{3\tau_s} \zeta^3 - \frac{1}{\tau_s^2} + \frac{20}{3\tau_s} - \frac{721}{45} + \frac{896}{15} \tau_s - \frac{1637}{15} \tau_s^2 \\ &\quad \left. + 36\tau_s^3 + \tau_s^4 - \frac{(1-\tau_s)^7}{\tau_s^{\frac{3}{2}}} (\operatorname{arc} \tanh \zeta \tau_s^{\frac{1}{2}} - \operatorname{arc} \tanh \tau_s^{\frac{1}{2}}) \right\}. \end{aligned}$$

Moreover, the expansion of $e^{\nu\chi_\nu(\tau)}$ equivalent to (2.9) is valid uniformly in $0 \leq \tau \leq \tau_s - \epsilon$ and $|\arg \nu| \leq \pi - \epsilon$.

This is proved from (2.6) by the saddle-point method; except that the explicit calculation of the coefficients p_r by this method would be very laborious, and the values shown in (2.13) were found by formal solution of the differential equation (1.2) in the form (2.9).

We begin by inserting in (2.6) the expansions of a_ν , b_ν and an asymptotic expansion (resembling (2.3)) of the gamma-function factor, so that the integrand takes the form

$$\exp \{ \nu f(t) + f_0(t) + f_1(t)/\nu + \dots \},$$

$$f(t) = \alpha \log \alpha - (1+\alpha) \log (1+\alpha) + \frac{1}{2} \log \tau + (1+\alpha) \log \frac{1-t}{1-\tau t} - \alpha \log (-t). \quad (2.14)$$

The saddle-points of $f(t)$ are at

$$t = \frac{\tau \tau_s^{-\frac{1}{2}} - 1 \pm (1-\tau)^{\frac{1}{2}} (1-\tau/\tau_s)^{\frac{1}{2}}}{2\alpha\tau} = t_1, t_2, \text{ say.} \quad (2.15)$$

For $0 < \tau < \tau_s$, t_1, t_2 are both real-negative, and the 'surface' of $\operatorname{Re} f(t)$ is as follows:† Going from right to left along the real t -axis we have a top at $t = \tau^{-1}$, a bottom at $t = 1$, a top at $t = 0$; from this top we fall to the saddle t_1 , then rise to the saddle t_2 , and then fall to $-\infty$. From t_1 valleys descend round the flanks of the top at 0 to the bottom at 1, and for each of these valleys the outer rim is a ridge ascending from

† See figure 1 (i). For $\tau \rightarrow 0$, $t_2 \rightarrow -\infty$ but $t_1 \rightarrow -\alpha$. It is because t_1 has a finite limiting position that we get the uniformity down to $\tau = 0$ stated at the end of the preceding enunciation.

t_2 to the top τ^{-1} . Thus for the integral in (2.6) the path can be taken to descend from the saddle t_1 to the bottom 1 through the two valleys, and the asymptotic expansion of $\chi_\nu(\tau)$ is found, by routine calculations, as the contribution to the integral from the neighbourhood of t_1 . This yields the formula (2.9).

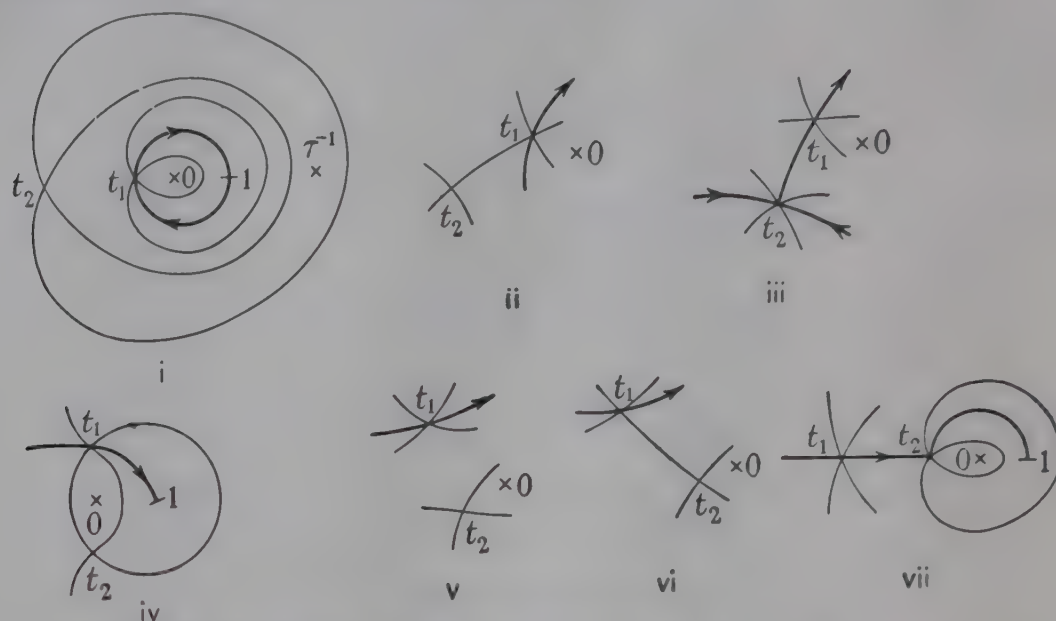


FIGURE 1. Contours $\text{Rl } f(t) = \text{const.}$, and path of integration, in transitional cases. The values of $\arg(\tau_s - \tau)$ are (i) 0, (ii) $\frac{1}{3}\pi$ (approx., for $|\tau_s - \tau|$ small), (iii) $\frac{2}{3}\pi$ (approx.), (iv) π , (v) $\frac{4}{3}\pi$ (approx.), (vi) $\frac{5}{3}\pi$ (approx.), (vii) 2π .

By taking the steepest path down from t_1 , on each side, the formula is established for $|\arg \nu| < \frac{1}{2}\pi$. The extension to $|\arg \nu| < \pi$ is made by using an 'isogonal' path which cuts the contours $\text{Rl } f(t) = \text{const.}$ at a constant acute angle (in one or the other sense according as $\text{Im } \nu$ is positive or negative) instead of at right angles. Such an isogonal path approaches $t = 1$ spirally, as noted after the enunciation of theorem 2.

The asymptotic expansions of $\frac{1}{2}\chi h_\nu^{(1)}(\tau)$, $\frac{1}{2}\chi h_\nu^{(2)}(\tau)$ are similarly found from the integrals in (2.7), using paths of integration which pass through the saddle t_2 as in figure 1 (vii). It will be shown below how these expansions can be deduced, without further calculation, from (2.9).

THEOREM 4. *Let the analytic continuations of the coefficients in the series on the right of (2.9) be taken, for τ circling in the positive sense round τ_s , starting on the segment $0 < \tau < \tau_s$, crossing $\tau_s < \tau < 1$ and returning to $0 < \tau < \tau_s$. Then at the half-way stage the series gives the asymptotic expansion of $\frac{1}{2}\chi h_\nu^{(1)}(\tau)$, valid for $\tau_s < \tau < 1$ and $|\arg \nu| < \pi$, and at the final stage it gives the asymptotic expansion of $\frac{1}{2}\chi h_\nu^{(1)}(\tau)$, valid for $0 < \tau < \tau_s$ and $|\arg \nu| < \frac{1}{2}\pi$.*

If τ circles in the negative sense, the asymptotic expansions obtained are those of $\frac{1}{2}\chi h_\nu^{(2)}(\tau)$.

Proof. The positions of the saddle t_1 , and of the steepest path descending from it (so far as the neighbourhood of t_1 is concerned), vary continuously with τ ; and for integration along the steepest path, the contribution to $\int \phi(t, \tau, \nu) dt$, from the neighbourhood of t_1 , is obtained by analytic continuation from the series on the right of (2.9). For any chosen τ let the ends of the steepest path be at A, B (each of these is

necessarily one of the bottoms $1, \infty$). Then, since t_1 is the highest point on the path, the analytically continued series (2.9) will give the asymptotic expansion of

$$\int_A^B \phi(t, \tau, \nu) dt.$$

When τ circles in the positive sense round τ_s the most steeply descending path from the saddle t_1 varies as shown in figure 1. Its upper limit B remains at $t = 1$, but the lower limit A switches† from 1 to ∞ at stage (iii), and thereafter remains at ∞ . After the switch the path is equivalent to $(i\infty, 1)$, so by (2.7) and (2.8) the continued series (2.9) represents $\frac{1}{2}\chi h_\nu^{(1)}(\tau)$.

We see further that at the half-way stage (iv) an isogonal path through t_1 , inclined at any acute angle (in either sense) to the contours $\text{Rl}f(t) = \text{const.}$, is equivalent to $(i\infty, 1)$, so the range of validity of the asymptotic expansion is $|\arg \nu| < \pi$. But at the final stage (vii), an isogonal path starting from t_1 into the lower half-plane is equivalent to $(-i\infty, 1)$, and the range in which the expansion represents $\frac{1}{2}\chi h_\nu^{(1)}(\tau)$ is $-\pi < \arg \nu < \frac{1}{2}\pi$. For systematic purposes the restriction $|\arg \nu| < \frac{1}{2}\pi$ of the enunciation is sufficient.

COROLLARY 1. If τ circles twice round τ_s in the positive sense, the final position of the steepest path through t_1 is the same as its initial position but with reversed sense, so the continuation of the series (2.9) must be minus the original series. The factor $(1 - \tau/\tau_s)^{-\frac{1}{2}}$ accounts for the change in sign; hence $u, p_1, p_2, \dots$ must be functions of τ with simple branch-points at τ_s . This is verified in the formulae given for $u, p_1, \dots, p_4$.

COROLLARY 2. For $|\arg(\tau_s - \tau)| \leq \pi$ let the ‘principal branch’ of each of $u, p_1, \dots$ be defined by taking in (2.10) to (2.13) the principal values of the square roots (so that u, ζ are real-positive when $0 < \tau < \tau_s$). Then the ‘subsidiary branch’ is obtained by reversing the signs of u, ζ ; let it be denoted by $\bar{u}, \bar{p}_1, \dots$.

Now let $0 < \tau < \tau_s$, and let $\arg \nu$ be in the neighbourhood of $\frac{1}{4}\pi$. Then by theorems 3 and 4,

$$\chi_\nu(\tau) \sim \frac{(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}}{(1-\tau/\tau_s)^{\frac{1}{2}}} \delta^\nu e^{-\nu u} \left(1 + \frac{p_1}{\nu} + \dots\right), \quad (2.16)$$

$$\frac{\chi_{-\nu}(\tau)}{2i\pi h_\nu} - \frac{e^{-\nu\pi i}\chi_\nu(\tau)}{2i\sin\nu\pi} = \frac{1}{2}\chi h_\nu^{(1)}(\tau) \sim \frac{(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}}{i(1-\tau/\tau_s)^{\frac{1}{2}}} \delta^\nu e^{-\nu\bar{u}} \left(1 + \frac{\bar{p}_1}{\nu} + \dots\right); \quad (2.17)$$

and since $u > 0$ and $\bar{u} = -u < 0$, the second term on the left of (2.17) is exponentially small compared with the one on the right, and can be omitted. Since (2.16) is valid for $|\arg \nu| < \pi$ we can here replace ν by $-\nu$, and by comparison with (2.17) there follows

$$1 + \sum_1^\infty p_r (-\nu)^{-r} \sim 2\pi h_\nu \delta^{2\nu} \left(1 + \sum_1^\infty \bar{p}_r \nu^{-r}\right). \quad (2.18)$$

Substitute here the expansion of $2\pi h_\nu \delta^{2\nu}$ obtained by taking the exponential of (2.3), and we obtain a formal power-series identity; and thence we can return to (2.18) under the sole restriction that the expansion (2.3) be valid, viz. for $|\arg \nu| < \pi$.

† The switch does not affect the continuity of the asymptotic expansion of the integral, since it affects only a part of the path which is lower than t_1 .

From theorems 3, 4 along with (2.18) it is easy to deduce:

THEOREM 5. For $0 < \tau < \tau_s$ and $|\arg \nu| < \frac{1}{2}\pi$,

$$(i) \chi y_\nu(\tau) \sim -\frac{2(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}}{(1-\tau/\tau_s)^{\frac{1}{2}}} \delta^\nu e^{\nu u} \left(1 + \frac{\bar{p}_1}{\nu} + \frac{\bar{p}_2}{\nu^2} + \dots\right), \quad (2.19)$$

$$(ii) \chi_{-\nu}(\tau) \sim \frac{(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}}{(1-\tau/\tau_s)^{\frac{1}{2}}} \delta^{-\nu} \left[e^{\nu u} \left(1 - \frac{p_1}{\nu} + \frac{p_2}{\nu^2} - \dots\right) + \frac{1}{2} \cot \nu \pi e^{-\nu u} \left(1 - \frac{\bar{p}_1}{\nu} + \frac{\bar{p}_2}{\nu^2} - \dots\right) \right]. \quad (2.20)$$

(iii) Formula (2.9) is valid without restriction upon $\arg \nu$, provided that, for every negative integer $-n$, $|\nu + n| > e^{-2n(u-c)}$, where c is a positive constant; i.e. (roughly speaking) ν is not near any negative integer.

$$(iv) \chi_{-n}^*(\tau) \sim \frac{(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}}{(1-\tau/\tau_s)^{\frac{1}{2}}} \delta^{-n} e^{nu} \left(1 - \frac{p_1}{n} + \frac{p_2}{n^2} - \dots\right). \quad (2.21)$$

THEOREM 6. For $\tau_s < \tau < 1$ let

$$\omega = \tau_s^{-\frac{1}{2}} \arctan \sqrt{\frac{\tau - \tau_s}{1 - \tau}} - \arctan \sqrt{\frac{\tau/\tau_s - 1}{1 - \tau}}, \quad > 0, \quad (2.22)$$

let P_r denote the continuation through the lower half-plane of p_r (as defined by (2.13) for $0 < \tau < \tau_s$), and let $\bar{P}_r$ be its conjugate. Then for $|\arg \nu| < \pi$,

$$\chi_\nu(\tau) \sim \frac{(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}}{(\tau/\tau_s - 1)^{\frac{1}{2}}} \delta^\nu \left[e^{i(\nu\omega - \frac{1}{2}\pi)} \left(1 + \frac{P_1}{\nu} + \frac{P_2}{\nu^2} + \dots\right) + e^{-i(\nu\omega - \frac{1}{2}\pi)} \left(1 + \frac{\bar{P}_1}{\nu} + \frac{\bar{P}_2}{\nu^2} + \dots\right) \right], \quad (2.23)$$

$$\chi y_\nu(\tau) \sim \frac{(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}}{(\tau/\tau_s - 1)^{\frac{1}{2}}} \delta^\nu \left[-i e^{i(\nu\omega - \frac{1}{2}\pi)} \left(1 + \frac{P_1}{\nu} + \dots\right) + i e^{-i(\nu\omega - \frac{1}{2}\pi)} \left(1 + \frac{\bar{P}_1}{\nu} + \dots\right) \right], \quad (2.24)$$

$$\chi_{-\nu}(\tau) \sim \frac{(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}}{(\tau/\tau_s - 1)^{\frac{1}{2}}} \delta^{-\nu} \left[\frac{i e^{-i\nu\pi - i(\nu\omega + \frac{1}{2}\pi)}}{2 \sin \nu \pi} \left(1 - \frac{P_1}{\nu} + \frac{P_2}{\nu^2} - \dots\right) - \frac{i e^{i\nu\pi + i(\nu\omega + \frac{1}{2}\pi)}}{2 \sin \nu \pi} \left(1 - \frac{\bar{P}_1}{\nu} + \frac{\bar{P}_2}{\nu^2} - \dots\right) \right], \quad (2.25)$$

$$\chi_{-n}^*(\tau) \sim \frac{(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}}{(\tau/\tau_s - 1)^{\frac{1}{2}}} \delta^{-n} \left[e^{-i(n\omega + \frac{1}{2}\pi)} \left(1 + \frac{i h'_n}{\pi h_n}\right) \left(1 - \frac{P_1}{n} + \frac{P_2}{n^2} - \dots\right) + e^{i(n\omega + \frac{1}{2}\pi)} \left(1 - \frac{i h'_n}{\pi h_n}\right) \left(1 - \frac{\bar{P}_1}{n} + \frac{\bar{P}_2}{n^2} - \dots\right) \right], \quad (2.26)$$

where by (2.3) $h'_n/h_n \sim -2 \log \delta - c_1 n^{-2} - 3c_3 n^{-4} - \dots$

Asymptotic series for the ψ -functions. These have the form of those given for the χ -functions, with $(1-\tau)^{\frac{1}{2}-\frac{1}{2}\beta}$ replaced by $(1-\tau)^{\frac{1}{2}+\frac{1}{2}\beta}$, and $p_1, p_2, \dots$ replaced by $\tilde{p}_1, \tilde{p}_2, \dots$. The latter are given by equations like (2.13) with ϕ replaced by $\tilde{\phi}$, where

$$\tilde{\phi} = \phi - \frac{2(1-\zeta^2)(1-\tau_s \zeta^2)}{(1-\tau_s) \zeta^2}.$$

In particular

$$\left. \begin{aligned} \tilde{p}_1 &= p_1 + 1 - \zeta, \\ \tilde{p}_2 &= \frac{1}{2} \tilde{p}_1^2 + \frac{1}{4} \tilde{\phi}, \\ \tilde{p}_3 &= p_3 + p_2(\zeta - p_1) + \tilde{p}_2(1 + p_1) - p_1 \tilde{p}_1. \end{aligned} \right\} \quad (2.27)$$

Numerical details. We must expect that, when any of the preceding series is curtailed, the error will be indicated in the usual way by the magnitude of the first neglected term.† Hence it is of primary interest to examine how the coefficients p_r grow as τ approaches τ_s , and this is shown by the specimen values given in table 1. From this table it is clear that there is a wide range surrounding τ_s in which the series will give, at best, very crude approximations unless ν is prohibitively large; for example, if we require values correct to 3 significant figures, then for $\nu = \pm 10$ we should exclude τ from about the range (0.07, 0.34). In the next section series free from this defect are given.

TABLE 1. SPECIMEN COEFFICIENTS IN ELEMENTARY ASYMPTOTIC SERIES (§ 2)

(Computed from formulae (2.13), (2.27) for $\gamma = 1.405^*$ giving $\tau_s = 0.168399$.)

τ	p_1	p_2	p_3	$\tilde{p}_1$	$\tilde{p}_2$	$\tilde{p}_3$
0.02	-0.05531	0.06705	-0.0807	-0.00358	0.00835	-0.0057
0.04	-0.12673	0.1928	-0.3379	-0.01793	0.05004	-0.1002
0.06	-0.22475	0.4496	-1.178	-0.05227	0.1809	-0.5760
0.08	-0.37110	1.0404	-4.203	-0.12647	0.5705	-2.720
0.10	-0.61802	2.665	-17.84	-0.28981	1.841	-13.87
0.12	-1.12372	8.666	-108.6	-0.69521	7.104	-95.6
0.14	-2.62143	48.69	-1493.1	-2.06426	45.00	-1422.5
0.20	-0.27735	-51.26	13.22	0.72265	-49.04	-36.23
	-2.41594i	+0.67i	+1925.5i	-2.90026i	-2.10i	+1894.9i
0.24	-0.27735	-5.628	1.646	0.72265	-4.302	-2.823
	-0.60249i	+0.167i	+65.621i	-1.35045i	-0.976i	+60.014i
0.30	-0.27735	-1.317	0.450	0.72265	-0.3735	0.016
	-0.07097i	+0.020i	+6.483i	-1.12757i	-0.8148i	+4.179i
0.40	-0.27735	-0.4595	0.213	0.72265	0.1390	0.386
	+0.21469i	-0.0595i	+0.892i	-1.29930i	-0.9389i	-0.886i
0.50	-0.27735	-0.2845	0.164	0.72265	-0.0807	0.227
	+0.35093i	-0.0973i	+0.307i	-1.63358i	-1.1805i	-1.773i

* This table was computed at a time when it seemed preferable to take $\gamma = 1.405$ rather than $\gamma = 1.4$. For practical purposes it is superseded by table 2, so it has not been recomputed with $\gamma = 1.4$.

3. UNIFORM ASYMPTOTIC FORMULAE

In an approximation to $\chi_\nu(\tau)$ which is uniform for τ near τ_s and ν large, it is necessary that the leading term should imitate $\chi_\nu(\tau)$ in being monotonic for $\tau < \tau_s$ but rapidly oscillating for $\tau > \tau_s$. Functions of this character may be obtained as solutions of differential equations which resemble (1.2) in the essential feature that the coefficient of $\nu^2 y$ has a simple zero at $\tau = \tau_s$; and when a companion differential equation has been chosen, its solutions (for $\nu \sim \infty$) can be compared with those of (1.2) by a method which has been expounded particularly by Langer (1931). Indeed, it is possible to determine complete asymptotic expansions for the solutions of (1.2), uniformly valid for τ near τ_s , in which the leading terms are properly chosen solutions of the companion equation.

† This statement, and those that follow, are intended only at the level of 'common-sense' judgement.

For practical purposes it is necessary that the solutions of the companion equation be tabulated functions. This condition is satisfied by the Bessel function $J_\nu(\nu t)$, whose differential equation is

$$\frac{d^2 z}{dt^2} + \frac{1}{t} \frac{dz}{dt} + \nu^2 \left(1 - \frac{1}{t^2}\right) z = 0; \quad (3.1)$$

and this is indicated as a suitable companion to (1.2) since in both cases the values of ν for which logarithmic solutions occur are the integers. The appropriate choice of t as a function of τ is found as follows:

In (1.2) take u , defined by (2.10), as independent variable, and put

$$w = \frac{(1 - \tau/\tau_s)^{\frac{1}{2}} y}{(1 - \tau)^{\frac{1}{2} - \frac{1}{2}\beta}}, \quad [y = \chi_\nu(\tau)]; \quad (3.2)$$

we obtain
$$\frac{d^2 w}{du^2} = (\nu^2 - \phi) w, \quad (3.3)$$

where ϕ is defined by (2.12) and (2.11); and near $\tau = \tau_s$, where $u = 0$, there is a development of the form

$$\phi = \frac{5}{36u^2} + \frac{A}{u^{\frac{1}{2}}} + A_0 + A_1 u^{\frac{1}{2}} + \dots$$

In (31) take the variables

$$u = \operatorname{arc tanh} \sqrt{(1 - t^2)} - \sqrt{(1 - t^2)}, \quad W = (1 - t^2)^{\frac{1}{2}}, \quad [z = J_\nu(\nu t)], \quad (3.4)$$

and we obtain
$$\frac{d^2 W}{du^2} = (\nu^2 - \Phi) W, \quad (3.5)$$

where
$$\Phi = \frac{t^2(1 + \frac{1}{4}t^2)}{(1 - t^2)^3}; \quad (3.6)$$

and near $t = 1$, where $u = 0$, there is a development

$$\Phi = \frac{5}{36u^2} + \frac{B}{u^{\frac{1}{2}}} + B_0 + B_1 u^{\frac{1}{2}} + \dots$$

Comparing (3.3) with (3.5), it is seen that when ν is large, $(\nu^2 - \phi)/(\nu^2 - \Phi)$ is near to 1 in a range of u surrounding $u = 0$. Hence it is appropriate so to relate t to τ that (as anticipated by the notation) (2.10) and (3.4) define the same u , i.e.

$$\operatorname{arc tanh} \sqrt{(1 - t^2)} - \sqrt{(1 - t^2)} = u = \operatorname{arc tanh} \sqrt{\frac{1 - \tau/\tau_s}{1 - \tau}} - \sqrt{\frac{1 - \tau/\tau_s}{1 - \tau}} \operatorname{arc tanh} \sqrt{\frac{\tau_s - \tau}{1 - \tau}}. \quad (3.7)$$

This relation gives t^2 as a function of τ without singularity at $\tau = \tau_s$ or $\tau = 0$. For τ near τ_s , t^2 is near 1; and the expanded form of (3.7) is

$$(1 - t^2)^{\frac{3}{2}} + \frac{3}{5}(1 - t^2)^{\frac{5}{2}} + \dots = (1 - \tau_s) \left\{ \left(\frac{1 - \tau/\tau_s}{1 - \tau} \right)^{\frac{3}{2}} + \frac{3(1 + \tau_s)}{5} \left(\frac{1 - \tau/\tau_s}{1 - \tau} \right)^{\frac{5}{2}} + \dots \right\}, \quad (3.8)$$

whence
$$1 - t^2 = (1 - \tau_s)^{\frac{2}{3}} \left(\frac{1 - \tau/\tau_s}{1 - \tau} \right) + O\left(\frac{1 - \tau/\tau_s}{1 - \tau} \right)^2.$$

For τ near 0 we find the power-series development

$$t^2 = \frac{4\tau}{e^2 \delta^2} + O(\tau^2), \quad (3.9)$$

and for $\tau_s < \tau < 1$ the continuation of (3.7) is

$$\sqrt{(t^2 - 1)} - \arctan \sqrt{(t^2 - 1)} = \omega = \tau_s^{-\frac{1}{2}} \arctan \sqrt{\frac{\tau - \tau_s}{1 - \tau}} - \arctan \sqrt{\frac{\tau/\tau_s - 1}{1 - \tau}}. \quad (3.10)$$

Starting from the differential equations (3.3), (3.5) the author (Cherry 1950) has proved the following:

THEOREM 7. For $\nu \sim \infty$ with $|\arg \nu| \leq \pi - \epsilon$ and $0 \leq \tau \leq 1 - \epsilon$, uniformly:

$$(i) \quad \chi_\nu(\tau) \sim N \left(\frac{\nu}{h_\nu} \right)^{\frac{1}{2}} \left[J_\nu(\nu t) \left(1 + \frac{q_2}{\nu^2} + \frac{q_4}{\nu^4} + \dots \right) + t J'_\nu(\nu t) \left(\frac{q_1}{\nu} + \frac{q_3}{\nu^3} + \dots \right) \right], \quad (3.11)$$

$$\frac{2\tau}{\nu} \chi'_\nu(\tau) \sim N \left(\frac{\nu}{h_\nu} \right)^{\frac{1}{2}} \left[J_\nu(\nu t) \left(\frac{q_1^*}{\nu} + \frac{q_3^*}{\nu^3} + \dots \right) + t J'_\nu(\nu t) \left(q_0^* + \frac{q_2^*}{\nu^2} + \frac{q_4^*}{\nu^4} + \dots \right) \right], \quad (3.12)$$

where t is related to τ by (3.7), the q_r, q_r^* are certain functions of τ , and

$$N = (1 - \tau)^{\frac{1}{2} - \frac{1}{2}\beta} \left(\frac{1 - t^2}{1 - \tau/\tau_s} \right)^{\frac{1}{2}}. \quad (3.13)$$

(ii) Similar formulae for $\chi y_\nu(\tau)$, $\chi h_\nu^{(1)}(\tau)$, $\chi h_\nu^{(2)}(\tau)$ and their derivatives are derived from (3.11), (3.12) by replacing the symbol J_ν by Y_ν , $H_\nu^{(1)}$, $H_\nu^{(2)}$ respectively.

$$(iii) \quad \chi_{-\nu}(\tau) \sim N \frac{\pi(\nu h_\nu)^{\frac{1}{2}}}{\sin \nu\pi} \left[J_{-\nu}(\nu t) \left(1 + \frac{q_2}{\nu^2} + \dots \right) + t J'_{-\nu}(\nu t) \left(\frac{q_1}{\nu} + \frac{q_3}{\nu^3} + \dots \right) \right], \quad (3.14)$$

$$- \frac{2\tau}{\nu} \chi'_{-\nu}(\tau) \sim N \frac{\pi(\nu h_\nu)^{\frac{1}{2}}}{\sin \nu\pi} \left[J_{-\nu}(\nu t) \left(\frac{q_1^*}{\nu} + \frac{q_3^*}{\nu^3} + \dots \right) + t J'_{-\nu}(\nu t) \left(q_0^* + \frac{q_2^*}{\nu^2} + \dots \right) \right]. \quad (3.15)$$

Note. In these formulae, $J'_\nu(\nu t) = dJ_\nu(\nu t)/d(\nu t)$, not $dJ_\nu(\nu t)/dt$.

The best way to calculate the q_r is as follows (the q_r^* will of course follow by formal differentiation of (3.11)). Let w , W , u be defined by (3.2), (3.4). Then

$$t(1 - t^2)^{\frac{1}{2}} J'_\nu(\nu t) = \frac{1}{\nu} \left(\frac{t^2 W}{2(1 - t^2)} - (1 - t^2)^{\frac{1}{2}} \frac{dW}{du} \right);$$

and (3.11) implies that the differential equation (3.3) is formally satisfied by

$$w = W \left(1 + \frac{Q_2}{\nu^2} + \frac{Q_4}{\nu^4} + \dots \right) - \frac{dW}{du} \left(\frac{Q_1}{\nu^2} + \frac{Q_3}{\nu^4} + \dots \right),$$

where

$$\left. \begin{aligned} Q_1 &= q_1(1 - t^2)^{\frac{1}{2}}, & Q_3 &= q_3(1 - t^2)^{\frac{1}{2}}, \dots \\ Q_2 &= q_2 + \frac{t^2 q_1}{2(1 - t^2)}, & Q_4 &= q_4 + \frac{t^2 q_3}{2(1 - t^2)}, \dots \end{aligned} \right\} \quad (3.16)$$

By using (3.5) to express $d^2 w/du^2$ in terms of W , dW/du , the conditions that (3.3) be formally satisfied come out as

$$\left. \begin{aligned} 2\dot{Q}_1 &= \phi - \Phi, & 2\dot{Q}_2 &= (\phi - \Phi) Q_1 + \ddot{Q}_1, \\ 2\dot{Q}_3 &= (\phi - \Phi) Q_2 + 2\Phi \dot{Q}_1 + Q_1 \dot{\Phi} + \ddot{Q}_2, & 2\dot{Q}_4 &= (\phi - \Phi) Q_3 + \ddot{Q}_3, \end{aligned} \right\} \quad (3.17)$$

where $\dot{Q}_1 = dQ_1/du$, etc. The constants which arise on integrating these are found from the limiting form of (3.11) for $\tau = t = 0$, which in virtue of (3.9) is found to be

$$\left(\frac{e}{\nu}\right)^\nu \Gamma(\nu+1) \left(\frac{2\pi h_\nu \delta^{2\nu}}{2\pi\nu}\right)^{\frac{1}{2}} \sim 1 + \frac{Q_1(0)}{\nu} + \frac{Q_2(0)}{\nu^2} + \dots;$$

we expand the left-hand member by means of (2.3) and Stirling's series, and equate coefficients to get the $Q_r(0)$.

From (3.17) and (3.16) we can find closed expressions for the Q_r, q_r in terms of τ and the related variable t , e.g.

$$q_1 = \frac{5}{24(1-t^2)^2} - \frac{1}{8(1-t^2)} + \frac{(1-\tau_s)^2}{8\tau_s^{\frac{3}{2}}(1-t^2)^{\frac{1}{2}}} \operatorname{arc tanh} \left(\frac{\tau_s - \tau}{1-\tau} \right)^{\frac{1}{2}} \\ - \frac{(1-t^2)^{-\frac{1}{2}}}{8(1-\tau_s)} \left[\frac{5}{3} \left(\frac{1-\tau}{1-\tau/\tau_s} \right)^{\frac{3}{2}} - (1+\tau_s) \left(\frac{1-\tau}{1-\tau/\tau_s} \right)^{\frac{1}{2}} + \frac{(1-2\tau_s)(1-4\tau_s)}{\tau_s} \left(\frac{1-\tau/\tau_s}{1-\tau} \right)^{\frac{1}{2}} \right]. \quad (3.18)$$

For τ near τ_s we can use (3.8) to express this as a power series in $(1-\tau/\tau_s)/(1-\tau)$; the negative powers cancel, and q_1 is regular at τ_s .

Similarly $q_2, q_3, \dots$ are regular at τ_s ; and this regularity checks that part of theorem 7 which asserts the uniformity of the approximations.

The explicit formulae have been worked out only in the case where $\gamma = 1.4$, giving $\beta = 2.5$, $\tau_s = \frac{1}{6}$. They are as follows:

$$\text{Let } \zeta = \left(\frac{1-6\tau}{1-\tau} \right)^{\frac{1}{2}}, \quad s = (1-t^2)^{\frac{1}{2}},$$

$$A = \frac{1}{2}\phi = \frac{9}{10\zeta^6} - \frac{63}{50\zeta^4} + \frac{39}{50\zeta^2} - \frac{19}{50} - \frac{1}{25}\zeta^2,$$

$$B = \frac{1}{2}\Phi = \frac{5}{8s^6} - \frac{3}{4s^4} + \frac{1}{8s^2},$$

$$C = \frac{3}{5\zeta^2} - \frac{1}{5} - \frac{2}{5}\zeta^2,$$

$$D = \frac{1}{8s} \int (\phi^2 - \Phi^2) du \\ = \frac{\zeta}{2000s} \left(-\frac{75}{\zeta^{10}} + \frac{315}{2\zeta^8} - \frac{4377}{20\zeta^6} + \frac{5981}{24\zeta^4} - \frac{9139}{48\zeta^2} + 208 + \frac{8}{3}\zeta^2 \right) \\ - \frac{625\sqrt{6}}{4608s} \operatorname{arc tanh} \frac{\zeta}{\sqrt{6}} + \frac{25}{1152s^{10}} - \frac{5}{128s^8} + \frac{11}{640s^6} - \frac{1}{384s^4},$$

$$E = \frac{\dot{\phi}}{8s} = \frac{\zeta}{1000s} \left(-\frac{1620}{\zeta^{10}} + \frac{3402}{\zeta^8} - \frac{2502}{\zeta^6} + \frac{798}{\zeta^4} - \frac{102}{\zeta^2} + 28 - 4\zeta^2 \right),$$

$$F = \dot{\Phi}/(8s) = \frac{1}{16} \left(-\frac{15}{s^{10}} + \frac{27}{s^8} - \frac{13}{s^6} + \frac{1}{s^4} \right),$$

$$J = \frac{1}{8}(\ddot{\phi} - \ddot{\Phi}) = -\frac{1}{16} \left(\frac{135}{s^{12}} - \frac{324}{s^{10}} + \frac{254}{s^8} - \frac{68}{s^6} + \frac{3}{s^4} \right) \\ + \frac{1}{2500} \left(\frac{43740}{\zeta^{12}} - \frac{122472}{\zeta^{10}} + \frac{128169}{\zeta^8} - \frac{62874}{\zeta^6} + \frac{14940}{\zeta^4} - \frac{1470}{\zeta^2} - 83 + 56\zeta^2 - 6\zeta^4 \right),$$

TABLE 2. COEFFICIENTS IN THE FORMULAE OF THEOREM 7

(Computed from formulae (3.19) for $\gamma = 1.4$, giving $\tau_s = \frac{1}{6}$.)

τ	$(1-\tau)^\beta$	t	q_1	q_2	q_3	q_4
0	1.0	0	0.36139 16	0.065302	-0.07237	-0.0283
0.02	0.95074 7494	0.33752 9219	0.36011 55	0.058741	-0.07126	-0.0279
0.04	0.90297 9899	0.47894 8474	0.35879 99	0.052239	-0.07007	-0.0276
0.06	0.85668 1984	0.58859 8619	0.35744 28	0.045803	-0.06879	-0.0275
0.08	0.81183 8360	0.68201 9208	0.35604 20	0.039439	-0.06742	-0.0277
0.10	0.76843 3472	0.76521 5279	0.35459 49	0.033158	-0.06596	-0.0281
0.12	0.72645 1593	0.84126 1828	0.35309 91	0.026968	-0.06439	-0.0287
0.14	0.68587 6824	0.91198 4274	0.35155 16	0.020880	-0.06273	-0.0295
0.16	0.64669 3082	0.97857 2336	0.34994 94	0.014904	-0.06094	-0.0305
0.18	0.60888 4097	1.04185 3617	0.34828 93	0.009055	-0.05903	-0.0317
0.20	0.57243 3402	1.10243 2428	0.34656 76	0.003344	-0.05699	-0.0334
0.22	0.53732 4331	1.16076 7056	0.34478 04	-0.002211	-0.05481	-0.0354
0.24	0.50354 0006	1.21721 5871	0.34292 34	-0.007595	-0.05247	-0.0378
0.26	0.47106 3332	1.27206 6395	0.34099 20	-0.012788	-0.04997	-0.0406
0.28	0.43987 6986	1.32555 4456	0.33898 09	-0.017768	-0.04728	-0.0439
0.30	0.40996 3413	1.37787 7288	0.33688 46	-0.022513	-0.04443	-0.0476
0.32	0.38130 4808	1.42920 2777	0.33469 69	-0.026996	-0.04135	-0.0519
0.34	0.35388 3113	1.47967 6168	0.33241 08	-0.031187	-0.03806	-0.0568
0.36	0.32768 0000	1.52942 5060	0.33001 87	-0.035052	-0.03452	-0.0623
0.38	0.30267 6863	1.57956 3208	0.32751 22	-0.038555	-0.03072	-0.0685
0.40	0.27885 4801	1.62719 3489	0.32488 18	-0.041652	-0.02662	-0.0755
0.42	0.25619 4607	1.67541 0263	0.32211 70	-0.044294	-0.02222	-0.0834
0.44	0.23467 6751	1.72330 1301	0.31920 57	-0.046426	-0.01747	-0.0923
0.46	0.21428 1363	1.77094 9410	0.31613 46	-0.047985	-0.01234	-0.1023
0.48	0.19498 8213	1.81843 3818	0.31288 85	-0.048898	-0.00678	-0.1135
0.50	0.17677 6695	1.86583 1427	0.30945 01	-0.049081	-0.00077	-0.1261

τ	q_0^*	q_1^*	q_2^*	q_3^*	q_4^*
0	1.0	0.36139 16	0.065302	-0.07237	-0.0283
0.02	1.00668 4299	0.36550 33	0.072489	-0.07403	-0.0289
0.04	1.01357 1340	0.37061 67	0.080029	-0.07589	-0.0294
0.06	1.02067 1755	0.37681 65	0.087941	-0.07796	-0.0298
0.08	1.02799 6981	0.38419 58	0.096246	-0.08023	-0.0300
0.10	1.03555 9332	0.39285 71	0.104968	-0.08270	-0.0301
0.12	1.04337 2093	0.40291 36	0.114131	-0.08537	-0.0298
0.14	1.05144 9623	0.41449 02	0.123761	-0.08822	-0.0291
0.16	1.05980 7460	0.42772 54	0.133887	-0.09123	-0.0279
0.18	1.06846 2459	0.44277 26	0.144539	-0.09440	-0.0263
0.20	1.07743 2935	0.45980 27	0.155750	-0.09768	-0.0244
0.22	1.08673 8825	0.47900 58	0.167555	-0.10105	-0.0220
0.24	1.09640 1891	0.50059 47	0.179993	-0.10447	-0.0192
0.26	1.10644 5926	0.52480 75	0.193104	-0.10787	-0.0156
0.28	1.11689 7013	0.55191 15	0.206933	-0.11119	-0.0111
0.30	1.12778 3818	0.58220 81	0.221527	-0.11433	-0.0057
0.32	1.13913 7925	0.61603 75	0.236938	-0.11719	+0.0008
0.34	1.15099 4231	0.65378 52	0.253221	-0.11964	0.0086
0.36	1.16339 1408	0.69588 99	0.270436	-0.12150	0.0178
0.38	1.17637 2444	0.74285 18	0.288648	-0.12257	0.0288
0.40	1.18998 5289	0.79524 42	0.307925	-0.12258	0.0417
0.42	1.20428 3608	0.85372 61	0.328342	-0.12120	0.0570
0.44	1.21932 7702	0.91905 93	0.349978	-0.11802	0.0751
0.46	1.23518 5587	0.99212 73	0.372917	-0.11255	0.0965
0.48	1.25193 4331	1.07396 08	0.397248	-0.10412	0.1217
0.50	1.26966 1652	1.16576 85	0.423063	-0.09194	0.1515

and let G, H be defined as below. Then

$$\begin{aligned}
 q_0^* &= \zeta/s, \\
 q_1 &= \frac{5}{24s^4} - \frac{1}{8s^2} - \frac{\zeta}{s} \left(\frac{1}{4\zeta^4} - \frac{7}{40\zeta^2} + \frac{1}{5} \right) + \frac{25\sqrt{6}}{48s} \operatorname{arc tanh} \frac{\zeta}{\sqrt{6}}, \\
 q_1^* &= \zeta s q_1 + C - \frac{\zeta}{2s} \left(\frac{1}{s^2} - 1 \right), \quad Q_2 = \frac{1}{2}(s^2 q_1^2 + A - B), \\
 q_2 &= Q_2 - \frac{1}{2} q_1 \left(\frac{1}{s^2} - 1 \right), \quad q_2^* = \frac{1}{2} \zeta s q_1^2 - \frac{\zeta}{2s} (A - B) + q_1 C, \\
 G &= q_1 Q_2 - \frac{1}{3} s^2 q_1^3 + D, \quad q_3 = G + q_1 B + E - F, \\
 q_3^* &= \zeta s (G - q_1 A - E + F) + Q_2 C - \frac{1}{2} q_2^* \left(\frac{1}{s^2} - 1 \right), \\
 H &= \frac{1}{2} s^2 q_1^2 (A - B) + A^2 - B^2 + 2s^2 q_1 (E + F) + J, \\
 Q_4 &= s^2 q_1 \{ q_3 - E + F - \frac{1}{2} q_1 (A + B) \} + \frac{1}{2} H - \frac{1}{4} Q_2 (s^2 q_1^2 - A + B), \\
 q_4 &= Q_4 - \frac{1}{2} q_3 \left(\frac{1}{s^2} - 1 \right), \quad q_4^* = \frac{\zeta}{s} (Q_4 - H) + q_3 C.
 \end{aligned} \tag{3.19}$$

From these formulae table 2 has been calculated; except that for $0.12 \leq \tau \leq 0.22$ series developments based on (3.8) had to be used (the exact formulae involve excessive loss of significant figures through cancellation).

The formulae of theorem 7, to the terms in ν^{-4} inclusive, give 7 significant figure accuracy for $|\nu| \geq 10$ and $0 \leq \tau \leq 0.5$. N may be found from table 2, since (3.13) gives

$$N = \{q_0^* (1 - \tau)^\beta\}^{-\frac{1}{\beta}}. \tag{3.20}$$

Also (2.3) becomes

$$\log (2\pi h_\nu)^{\frac{1}{2}} \sim 1.17344 \, 36154\nu + \frac{0.27805 \, 825}{\nu} - \frac{0.077461}{\nu^3} - \frac{0.0376}{\nu^5} + \dots, \tag{3.21}$$

and

$$\delta = 0.30929 \, 99950. \tag{3.22}$$

There are, of course, formulae for $\psi_\nu(\tau)$, etc., resembling those of theorem 7. It happens, however, that in these formulae the coefficients of ν^{-3} , ν^{-4} are of the order of 10 times the corresponding coefficients in (3.11); and numerical test verifies what would therefrom be guessed, that the ψ formulae are less accurate than the χ -formulae when curtailed at the same point. Hence to find $\psi_\nu(\tau)$ for large ν it is best to find χ_ν, χ'_ν from theorem 7 and then use (1.15).

Acknowledgement is made to Miss Betty Laby for much help in computing tables 1, 2.

REFERENCES

- Cherry, T. M. 1947 *Proc. Roy. Soc. A*, **192**, 45.
 Cherry, T. M. 1949 *Proc. Roy. Soc. A*, **196**, 1.
 Cherry, T. M. 1950 *Trans. Amer. Math. Soc.* **68**, 224.
 Garrick I. E. & Kaplan, C. 1944 *Rep. Nat. Adv. Comm. Aero., Wash.*, no. 790.
 Huckel, Vera 1948 *Tech. Notes Nat. Adv. Comm. Aero., Wash.*, no. 1716.
 Langer, R. E. 1931 *Trans. Amer. Math. Soc.* **33**, 23.
 Lighthill, M. J. 1947 *Proc. Roy. Soc. A*, **191**, 341.
 Seifert, H. 1947 *Math. Ann.* **120**, 75.
 Tsien, H. S. & Kuo, Y. H. 1946 *Tech. Notes Nat. Adv. Comm. Aero., Wash.*, no. 995.

Resistivity of pure metals at low temperatures

II. The alkaline earth metals

BY D. K. C. MACDONALD AND K. MENDELSSOHN

Clarendon Laboratory, University of Oxford

(Communicated by F. E. Simon, F.R.S.—Received 11 March 1950)

In continuation of the work reported in the first paper of this series the electrical resistivity of beryllium, magnesium, calcium, strontium and barium has been investigated in detail in the temperature region below 20° K.

In most cases the resistance was also determined in the range up to liquid-air temperatures in order to obtain estimates of the Debye characteristic temperatures, which are compared with data obtained by other methods.

A striking feature is the occurrence of a pronounced minimum in the resistance of magnesium below 10° K which was exhibited by all four specimens investigated.

1. INTRODUCTION

In the first paper in this series (MacDonald & Mendelssohn, 1950) we have presented a detailed investigation of the electrical conductivity of the alkali metals at low temperatures. This paper presents a similar examination of the alkaline earth metals. We have used this term in the wider sense, as is often done, to include the metals beryllium, and magnesium with calcium, strontium and barium. For obvious reasons, radium could not be investigated.

In theoretical approach all the alkali metals are generally treated as the ideal prototype exhibiting electrical conduction through quasi-free electrons. Even so it was found that only sodium showed close agreement throughout with the theory of Bloch & Grüneisen. The deviations observed—relatively strong in the case of the heavier alkalis—were interpreted as being owing to departures from the weak and isotropic electron binding assumed in the theory. It is evident that these deviations will be expected to become very much greater in the case of the alkaline earth metals since, moreover, the Brillouin zone structure must severely affect the electronic behaviour. In a simple divalent element electrical conductivity can only be realized if the first and second zones overlap partially, indicating further the complexity of the conduction process.

As is well known it is difficult to obtain the alkaline earths in a very pure form in the metallic state. We have naturally used in this work metals of the highest purity available but, except in the case of magnesium, the purity in general can probably not be regarded as being as high as that of the alkali metals used by us. We have assumed as before the validity of Matthiessen's rule for calcium, strontium and barium in this work; although the theoretical foundation for this is less certain than in monovalent metals. Beryllium and magnesium, as will be seen below, present special cases to which the rule does not seem to be applicable.

The experimental methods as regards measurement of electrical resistance, and the maintenance, control and measurement of the low temperatures were the same as in the previous work.

2. PREPARATION OF SPECIMENS

For calcium, strontium and barium thin strips were cut off blocks or lumps of these metals under sodium-dried (medical) liquid paraffin. Some magnesium wire of high purity was also available, and this was mounted directly on a small mica former. A number of attempts were made, as suggested by Bridgman (1927), to solder contacts with aluminium solder to this metal wire, but these were unsatisfactory, and it was found more profitable to make small screw boss contacts, first cleaning the surface of the wire under liquid paraffin. As also remarked by Bridgman, beryllium was found to present very considerable difficulty. The metal can be sawn, but great care is necessary to prevent the strip from curling up and cracking during the process.

The strips of these metals were then mounted in small ebonite blocks having a shallow, narrow groove cut in them through which small pointed screws to provide electrode contacts emerged slightly 'proud' from the surface. The metal was then compressed firmly between the mount and an opposing ebonite block, and usually a very satisfactory set of contacts was obtained which remained so at the lowest temperatures.

3. EXPERIMENTAL RESULTS

(i) *Beryllium*

Specimens from two different samples were examined. Analytical purities were not available, but they were stated to be of comparatively high purity.

The results for Be 1 and Be 2(a) are given first in table 1. In both cases it will be noticed that the resistance was effectively constant in the region below 20° K, and that both specimens displayed a relatively high residual resistance. Similar behaviour was found by Meissner & Voigt (1930) in two specimens of beryllium containing $\frac{1}{2}$ and 2 % Fe impurity respectively. It was remarkable, however, that in their work the more impure specimen exhibited a *lower* residual resistance (~ 0.307) than the purer one (~ 0.37). This behaviour is strongly suggestive of impurity semi-conductor properties, and concords with the fact (cf. Seitz 1940, p. 439) that the soft X-ray emission spectrum 'behaves as though the levels of a single zone were almost completely occupied, indicating that this metal is very nearly an insulator'.

TABLE 1. BERYLLIUM

T (° K)	Be 1		Be 2(a)	
	$R/R_{0^{\circ}\text{C}}$		$R/R_{0^{\circ}\text{C}}$	
273	1.00 ₀		1.00 ₀	
90	0.39 ₇		0.32 ₂	
20.4	0.38 ₄		0.27 ₆	
4.2	0.38 ₄		0.27 ₆	

Be 2(b). 'IDEAL' RESISTANCE (r)				
T (° K)	$r_{\text{obs.}}$	$\Theta = 650^{\circ}\text{K}$	$\Theta = 500^{\circ}\text{K}$	$\Theta = 700^{\circ}\text{K}$
		$r_{\text{calc.}}$	$r_{\text{calc.}}$	$r_{\text{calc.}}$
290	1.000	1.000	1.000	1.000
76.0	0.023 ₈	0.026 ₈	0.053 ₄	0.022 ₂
68.8	0.017 ₈	0.017 ₉	0.036 ₈	0.013 ₈
60.4	0.011 ₉	0.009 ₆	0.020 ₇	0.007 ₆
52.0	0.005 ₉	0.004 ₇	0.011 ₇	0.003 ₉
43.0	0.002 ₉	0.001 ₉	0.004 ₇	0.001 ₄

In view of the constancy of the resistance below 20°K , measurements were also made on another thinner specimen Be 2(b) from the second sample in the region between 90 and 20°K to derive an estimate of Θ by comparison with the Grüneisen law. It was found that $\Theta \approx 650^{\circ}\text{K}$ gave fair overall agreement, and for comparison, theoretical values for $\Theta \approx 500$ and 700°K are also given in table 1.

The value of Θ determined from specific heat measurements (Simon & Cristescu 1934) is $\sim 1000^{\circ}\text{K}$. It may be unjustifiable to consider the discrepancy as significant in view of the uncertainty in applying Grüneisen's formula to other than an ideal monovalent metal; on the other hand, it is interesting to observe that the ratio of the Θ values determined also for calcium (see below) is almost identical. Of the remaining alkaline earth metals it appears that specific heat data are available only for magnesium. In that case, however, the peculiar behaviour of the electrical resistance at low temperatures—to be discussed below—precludes a determination of Θ from Grüneisen's formula.

(ii) *Magnesium*

Four specimens in all were examined. Details of source and purity are:

specimen	source	purity (%)
Mg 1	Messrs Johnson Matthey (lab. no. 1703), wire drawn by them from their rod	> 99.95 (Mn: 0.03; Fe: 0.0075; Al: 0.004)
Mg 2	? ex Hopkins Williams	?
Mg 3	Messrs Johnson Matthey	second sample
Mg 4	Messrs New Metals and Chemicals Ltd.	Al 0.004 Zn 0.001 Ni < 0.0005 Ag < 0.0005 Cu < 0.001 Na < 0.001 Fe 0.001 Mn 0.002 Si 0.006 Ca 0.004 Pb 0.001 therefore, Mg > 99.978

All four specimens were remarkable in exhibiting a pronounced minimum in the curve of resistance against temperature in the region ~ 4 to 8°K (see Figure 1). This phenomenon in metallic conductors has so far only been established in the case of gold in the extensive measurements of van den Berg (1938). It is at present not at all clear how widespread is the phenomenon of resistance increase at very low temperatures. Van den Berg has observed a very small minimum in one of three specimens of silver, and Meissner & Voigt (1930) on occasion noted slightly higher resistances at the lowest helium temperatures for instance in sodium and aluminium. These observations have to be treated as isolated cases, since in subsequent experiments (for example, Boorse & Niewodniczański 1936) the effect was not reproduced. A decision on the reality of such very small effects, moreover, may be difficult if these are of the same order as perturbing thermal e.m.f.'s in the measuring circuit. It is therefore particularly significant that the minima observed in magnesium are of about the same relative size as those found in gold.

A direct experimental inference made by van den Berg was that the position of the minimum in temperature was a diminishing function of purity as measured by the residual resistance. It is clear from figure 1 that this general behaviour appears to be evinced for magnesium, although the range of purity available is too small to

justify quantitative inference of the functional relation. We should mention that Mg 4 was a rather thick specimen, and that the data therefore show more 'spread' than normal. Detailed location of the minimum, etc., in this case is therefore not warranted.

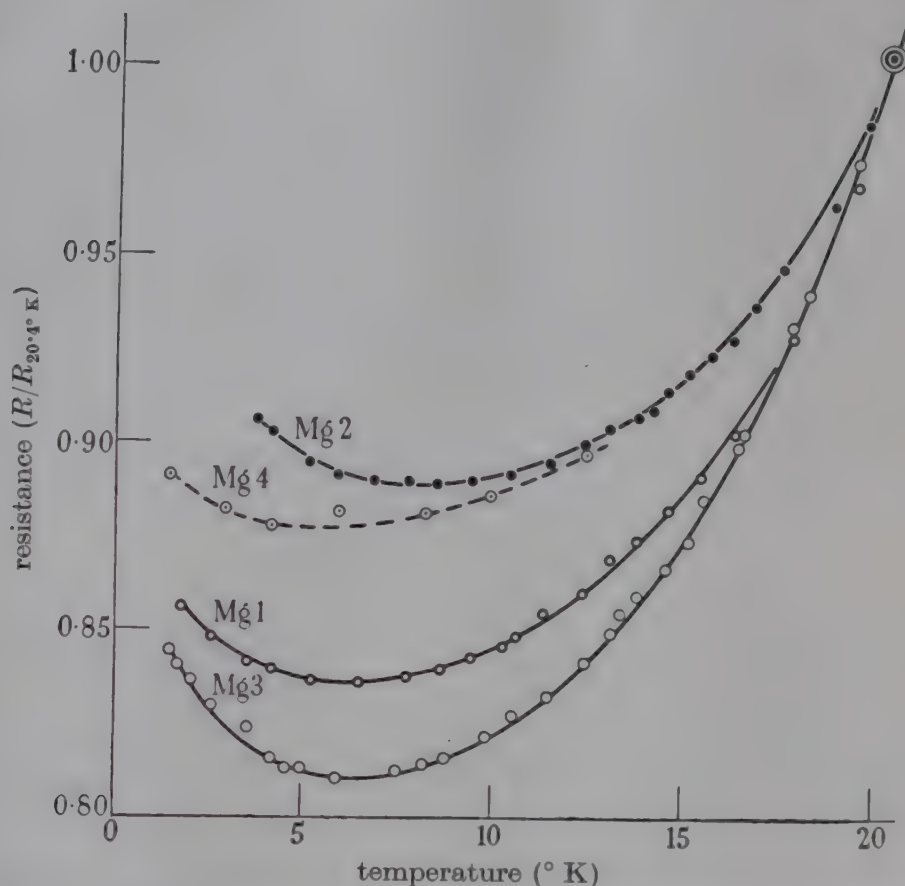


FIGURE 1. Temperature variation of electrical resistance of magnesium.

The physical cause for the minimum in gold has not yet been discovered. Justi (1948), discussing the question, has suggested that the geometrical size of the specimen may be responsible, but it seems difficult to sustain this idea. Variation of conductivity due to such an effect in a specimen of given size can arise from two sources: alteration of (1) the mean free path and (2) the scattering law at the surface. At most, if the former becomes great in comparison with the dimensions of the specimen, apparent *constancy* of the resistance can result; while any radical change of the latter, particularly at these very low temperatures, appears quite improbable.

Van der Leeden (1940) in his thesis suggested the inclusion of a term $A/\sqrt{T}$ in the general formula for the electrical resistance which would, of course, result in infinite resistance as $T \rightarrow 0$, but without theoretical justification. Very recently Lane (1949), in a letter to *Physical Review*, has recalled the problem but without any apparent fresh contribution to the subject. It has on occasion been suggested (cf. Gorter 1938) that quite generally metals will either become superconductive or tend to infinite resistance at the absolute zero, but it cannot be said that the available experimental evidence seriously indicates confirmation of his thesis. Furthermore, since the minimum descends in temperature with increasing purity, it is obvious that a highly peculiar situation—to say the least—would arise in this case for specimens of vanishing small impurity.

The only other case in metals where the resistance rises with falling temperature is in certain metals under the influence of a constant magnetic field, for instance in cadmium (cf. Milner 1937). The explanation there is that the magneto-resistive effect will, below a certain temperature, outweigh the thermal drop in resistivity. It seemed, interesting therefore, to see what would be the effect of an applied magnetic field in the case of magnesium. Results with fields of ~ 2 and ~ 4 kG are compared with the curve in zero field in figure 2. Although under a field of ~ 4 kG the resistance has been increased to almost double and the curve of its temperature dependence has been consequently flattened, the position of the minimum does not seem to have altered significantly. There are curious small deviations at ~ 2 and ~ 4 kG in the region below 5° K. While it would be premature to ascribe any definite significance to them it is possible that the true curve may inflect and become temperature-independent near the absolute zero. Experiments to elucidate this point at still lower temperatures are in progress.

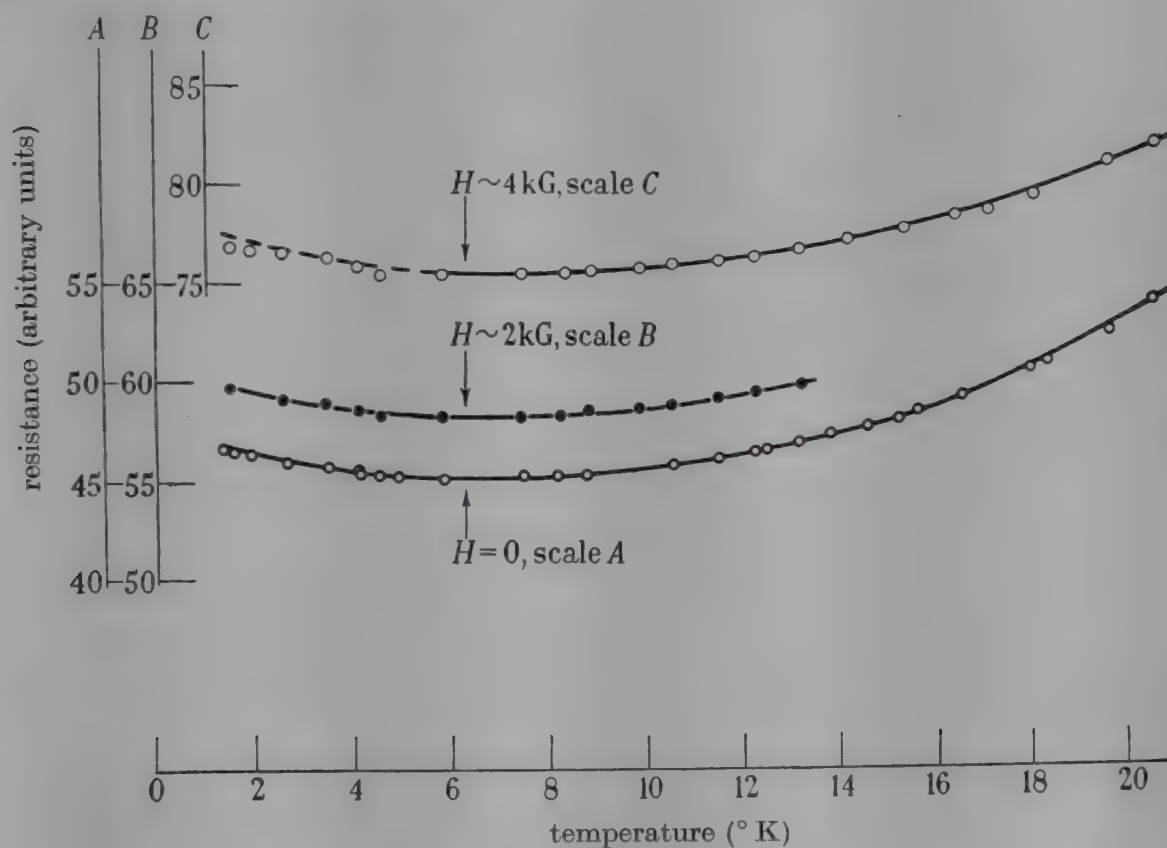


FIGURE 2. Magneto resistance of magnesium (specimen Mg 3).

TABLE 2. CALCIUM

T ($^\circ$ K)	$\Theta_{\text{exp.}}$ ($^\circ$ K)
80	144
40	150
20	137

(iii) *Calcium*

A sample of 'chemical purity' was obtained from Messrs Lights Ltd.

Resistance measurements were made continuously from 90° K downwards to $\sim 4^\circ$ K (see figure 3). The residual resistance is relatively high, rendering theoretical

comparison with the very low temperature data impossible. The detailed work between 80 and 20° K enabled relatively accurate determinations of Θ to be made in this region using Grüneisen's law. Representative values appear in table 2.

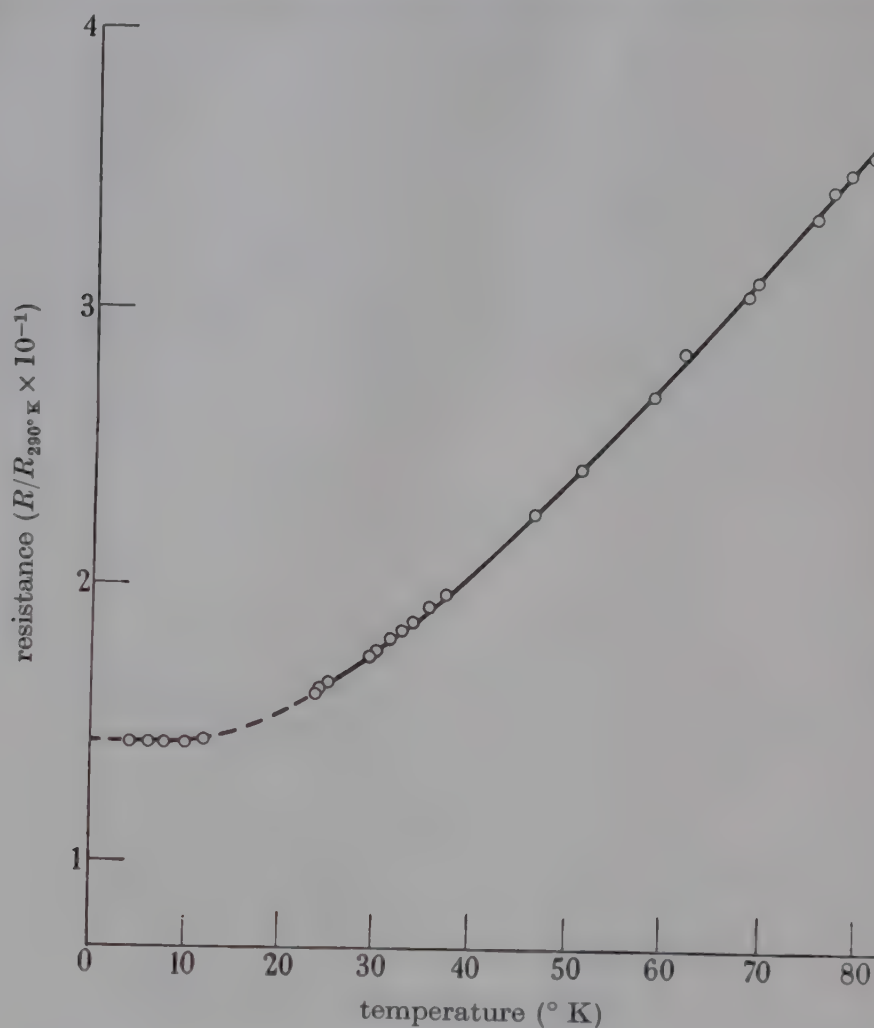


FIGURE 3. Electrical resistance of calcium.

We may remark that Meissner & Voigt reported that their specimen was too impure to permit any determination of Θ .

The value of Θ deduced by us (say 145° K mean) is considerably below that of 230° K obtained from specific heat measurements (cf. Simon 1931). We have noted above that the ratio (0.63) is practically the same as that found for beryllium.

(iv) *Strontium*

Two samples of 'chemical purity' from Messrs New Metals and Chemicals Ltd. and Messrs Lights Ltd. respectively were examined. The former was rather purer as judged by its residual resistance; attention is here confined to this specimen. In figure 4 the experimental data obtained below 80° K are presented and the curve below 20° K is shown inset in detail.

Grüneisen has deduced a value of $\Theta \sim 170^\circ \text{K}$ from electrical conductivities at the higher temperatures. Such a value, however, gives very poor agreement with our experimental results in the intermediate temperature range (~ 70 to 20°K), and we find that a value $\Theta \sim 100^\circ \text{K}$ gives much better agreement over this temperature

range (see table 3). We may note immediately that Lindemann's (1910) melting-point formula yields $\Theta \sim 136^\circ \text{K}$.^{*} Thus, we find a ratio of these values ~ 0.73 , substantially close to that found using specific heat values of Θ for beryllium and calcium (above).

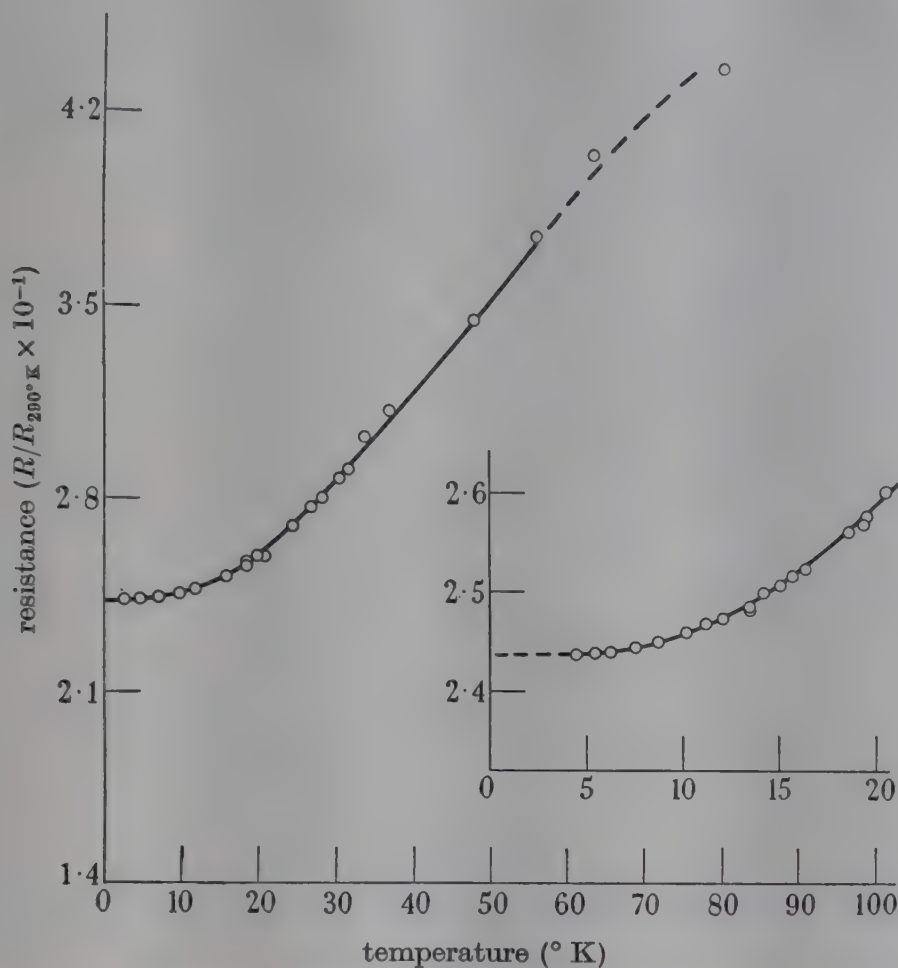


FIGURE 4. Electrical resistance of strontium.

TABLE 3. STRONTIUM, 'IDEAL' RESISTANCE (r)

T ($^\circ \text{K}$)	$r_{\text{obs.}}$	$\Theta = 170^\circ \text{K}$	$\Theta = 100^\circ \text{K}$
		$r_{\text{calc.}}$	$r_{\text{calc.}}$
290	1.000	1.000	1.000
63.3	0.212	0.151	0.194
55.9	0.173	—	0.164
48.2	0.132	0.090	0.134
37	0.090	—	0.088
34	0.076 ₆	—	0.076 ₂
28.7	0.049	0.022 ₂	0.054
20.4	0.021 ₂	—	0.024 ₃

Below $\sim 20^\circ \text{K}$ the divergence between experiment and the Grüneisen law becomes significant. This is evidenced by the logarithmic plot of figure 5. If the value $\Theta \sim 100^\circ \text{K}$ were valid overall we should expect a T^5 law to ensue below $\sim 12^\circ \text{K}$,

^{*} This is calculated using the same dimensionless parameter as derived for calcium by comparison with the specific heat value. The melting temperatures have been accepted by us as calcium, 851°C , and strontium, 771°C (Hoffmann & Schulze 1935).

while in fact an index $n \sim 3.4$ is observed in this region. The significance of deviations of this type at low temperatures will be discussed below.

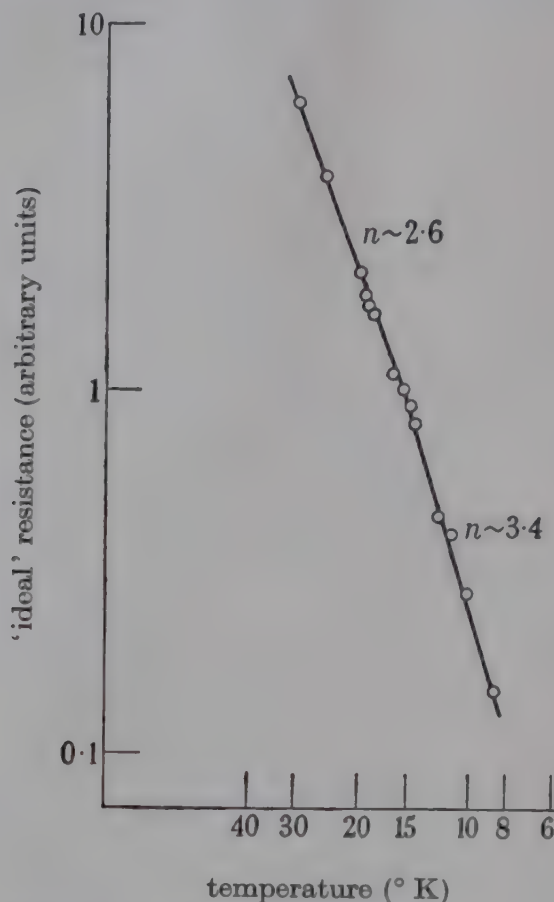


FIGURE 5. Logarithmic plot of variation of 'ideal' resistance of strontium.

(v) Barium

Metal was obtained through the courtesy of Mr L. F. Britten, of Messrs General Electric Co. Ltd., who estimated the purity to be in the order of 99 % with the lighter alkali metals as probable main impurities.

Figure 6 presents the curve obtained from $\sim 90^\circ \text{K}$ downwards, figure 7 the detailed results below 20°K , and figure 8 a logarithmic plot of the data. There seems to be little doubt that the resistance curve shows two small 'kinks'—at ~ 14 and $\sim 10^\circ \text{K}$ —of the same characteristic as those observed and discussed in the work on potassium and caesium published in our earlier paper. In view of the probable existence of sodium as a prime impurity here this observation seems of particular interest. As yet we can offer no detailed explanation of this remarkable behaviour in these metals except that in the investigations on a range of potassium specimens this type of anomaly was traced to the presence of a very small sodium impurity. The experimental data after subtraction of the residual resistance has been compared with the Grüneisen function, taking $\Theta = 100^\circ \text{K}$ ($\Theta = 104^\circ \text{K}^*$ deduced from the Lindemann melting-point formula).

It is evident that down to $\sim 20^\circ \text{K}$ the agreement is rather satisfactory, but severe divergence ensues rapidly thereafter in the same way as found in strontium.

* The melting-point has been taken as 704°C (Hoffmann & Schulze 1935).

TABLE 4. BARIUM—'IDEAL' RESISTANCE (r)

T ($^{\circ}$ K)	$r_{\text{obs.}}$	$\Theta = 100^{\circ}$ K $r_{\text{calc.}}$
290	1.000	1.000
82.5	0.252	0.264
60.0	0.171	0.180
40.0	0.103	0.100
20.3	0.024 ₆	0.023 ₄
10.7	0.0035 ₈	0.0021 ₈
8.21	0.00095 ₆	0.000068 ₇

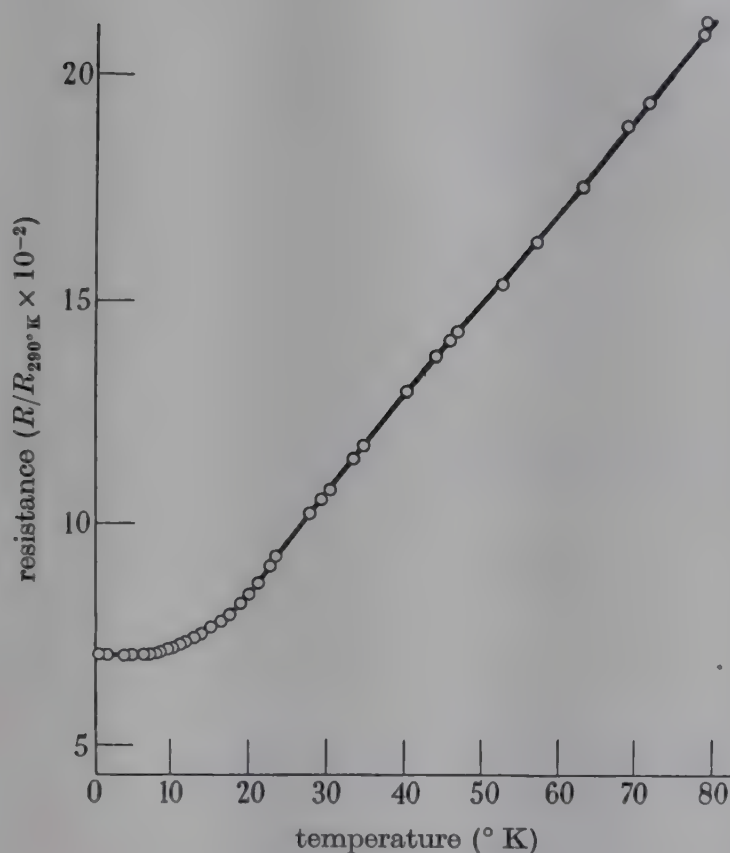


FIGURE 6. Electrical resistance of barium at low temperatures.

CONCLUSION

The most outstanding feature of the results is the minimum of resistivity in magnesium which has already been discussed.

Two main problems arise for discussion: first, how far we may attach significance to the Debye temperatures deduced from measurements in the intermediate temperature region; secondly, to what extent one may expect the individual electron-energy characteristics to play a part in the observed results. As pointed out in our first paper the entropy of a crystalline lattice only deviates appreciably from classical theory for $T < \sim 0.4\Theta$. We may therefore anticipate that the finer details of the electron-energy structure will only begin to make themselves felt at these lower temperatures—that is, in the case of strontium and barium particularly, more or less in the liquid-hydrogen and helium region. It is, then, to the complex electron-energy structure of the divalent metals that we ascribe the rather severe deviations from the Grüneisen law observed in the very low temperature region.

At intermediate temperatures, however, we should expect fair agreement with the universal Grüneisen law yielding a value for the characteristic temperature of real physical significance. While our deduced values for Θ in this region agree with those

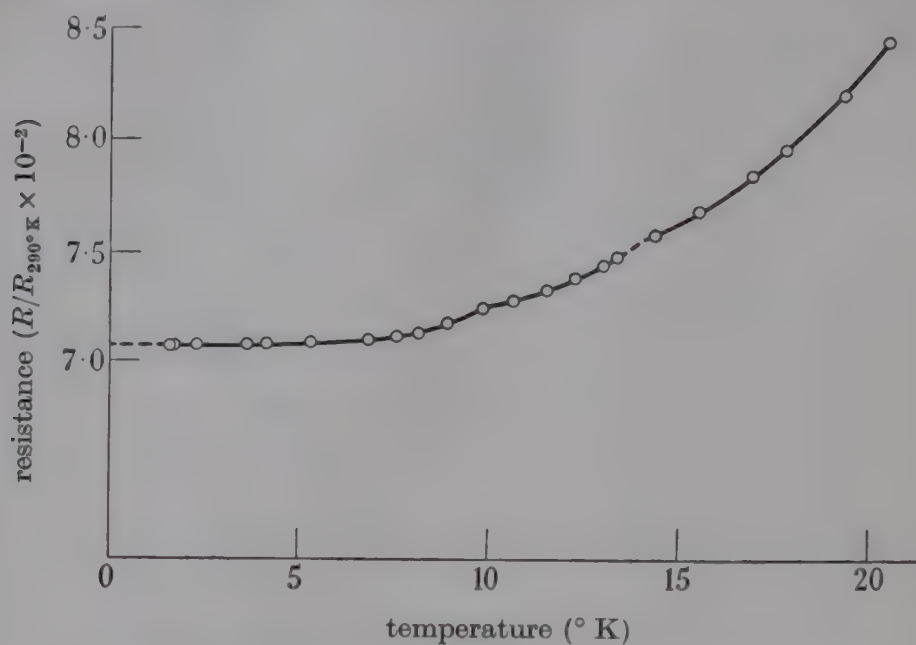


FIGURE 7. Electrical resistance of barium below 20° K.

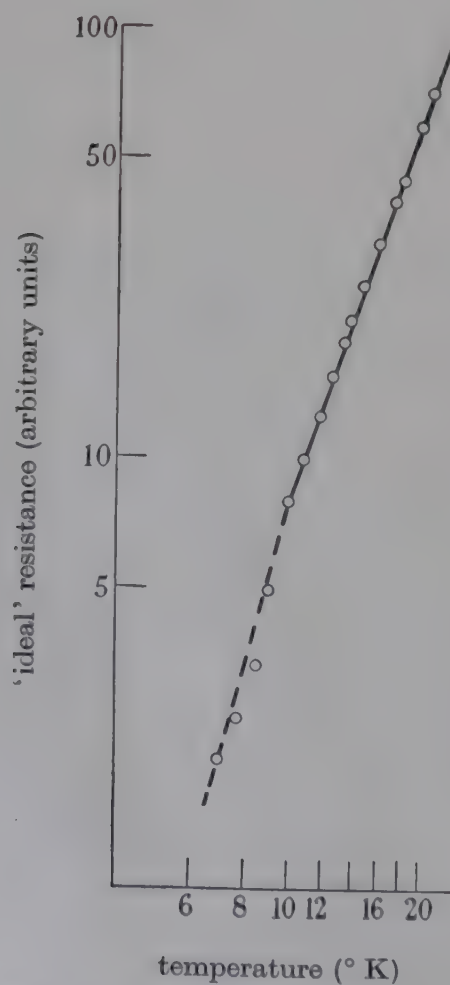


FIGURE 8. Logarithmic plot of 'ideal' resistance of barium below 20° K.

from other determinations in order of magnitude, this agreement is by no means close. It can clearly not be decided at this stage whether the existing discrepancies are due to the approximate nature of the Debye theory or to other causes, such as the possible excitation of internal degrees of freedom such as have been postulated for other metals (cf. Simon 1926, 1930).

We should like to thank Mr G. Williams for assistance with the measurements.

REFERENCES

- Berg, van den G. J. 1938 Thesis, Leiden.
Boorse, H. A. & Niewodniczański, H. 1936 *Proc. Roy. Soc. A*, **153**, 463.
Bridgman, P. W. 1927 *Proc. Amer. Acad. Arts Sci.* **62**, 207.
Gorter, C. J. 1938 *Physica*, **5**, 483.
Hoffmann, F. & Schulze, A. 1935 *Phys. Z.* **36**, 453.
Justi, E. 1948 *Leitfähigkeit und Leitungsmechanismus fester Stoffe*. Göttingen: Vandenhoeck und Ruprecht.
Lane, C. T. 1949 *Phys. Rev.* **76**, 304.
Leeden, van der 1940 Thesis, Leiden.
Lindemann, F. A. 1910 *Phys. Z.* **11**, 609.
MacDonald, D. K. C. & Mendelssohn, K. 1950 *Proc. Roy. Soc. A*, **202**, 103.
Meissner, W. & Voigt, B. 1930 *Ann. Phys., Lpz.*, **7**, 761, 892.
Milner, C. J. 1937 *Proc. Roy. Soc. A*, **160**, 207.
Seitz, F. 1940 *The modern theory of solids*. New York: McGraw-Hill.
Simon, F. & Cristescu, S. 1934 *Z. phys. Chem. (B)*, **25**, 273.
Simon, F. 1931 *Landolt-Börnstein Tabellen*, IIb, 1232.
Simon, F. 1926 *S.B. preuss. Akad. Wiss.* p. 477.
Simon, F. 1930 *Ergebn. exakt. Naturw.* **9**, 223.

Term values in hybrid states

By W. MOFFITT

British Rubber Producers' Research Association, Welwyn Garden City, Herts.

(Communicated by E. K. Rideal, F.R.S.—Received 14 March 1950)

In this paper a number of calculations are undertaken, which are considered necessary to a proper understanding of the factors governing orbital hybridization. A convenient classification of the different types of hybridization which arise for elements of the first short period is introduced. A general method is then developed for the calculation of the term values of the hybrid valence states associated with such atoms. This method is followed numerically for carbon, nitrogen and oxygen, which are particularly interesting in chemistry.

1. INTRODUCTION

The great success of Pauling & Slater's original theory (1931) of hybridization accounts, perhaps, for the relatively small number of fundamental studies by which their work has been succeeded. As the experimental methods of investigation have improved, however, details of molecular structure have been revealed with which the results of Pauling's theory of 1931 are inconsistent. The need for a more thorough discussion of the hybridization process has become apparent, and was recently stressed by Cottrell & Sutton (1947). In this and a following paper (Moffitt 1950), a number of calculations are undertaken which, it is hoped, may lead to a clearer understanding of the factors governing orbital hybridization.

Both for reasons of mathematical tractability and of physical intelligibility, it has been found convenient to refer the electronic structure of molecules to that of their constituent atoms. For some purposes this reference is most profitably made to atoms not in their spectroscopic states but in certain related valence states (cf. Van Vleck 1934). Such states are defined so that the intra-atomic couplings of the spins of their electrons correspond closely to those which may occur when the atoms become part of a molecule. The acceptance of this procedure commits one to the approximation of localization or of perfect electron pairing (in however general a sense). It is not intended to discuss the validity of the valence state concept here (cf. Voge 1936), but merely to stress an aspect of its nature which it has in common with the present work.

This first paper is largely concerned with the energies and ionization potentials of various valence states of the carbon, nitrogen and oxygen atoms. A knowledge of the energies of promotion required to produce different valence states of atoms from their ground states is of obvious thermochemical interest, and the importance of the ionization potentials of atoms in the analysis of molecular structures has been realized for some time. Thus a knowledge of the appropriate atomic term values may not only help to identify a molecular orbital whose ionization potential has been observed (Mulliken 1934*a*), but also to estimate the electronegativities of atoms under different circumstances (Mulliken 1934*b*). Interest in these atomic processes has recently been revived: on the empirical side, for example, Walsh (1946) has

related the observed properties of the carbonyl bond in a series of ketones to the ionization potential of the non-bonding electrons of the ketonic oxygen atom. And on the theoretical side, the term values of atomic valence states have been used as the basis of an approximate self-consistent field method for molecules (Moffitt 1949).

Mulliken (1934*b*) has related the ionization potentials of many atoms and ions in their simplest valence states to their observed spectra. But in the study of molecular structure many atomic valence states are encountered whose term values are less simply related to those of spectroscopic states. In these cases interest centres not on the ionization of, say, an *s* or *p* electron, but on that of an electron in an *s-p* hybrid orbital. Similarly, although Van Vleck (1934) has related the energies of valence states arising from the sp^3 configuration of carbon to the C I spectrum, the energies of promotion of states in 'second-order' hybridization (see below) remain to be calculated. Energies of promotion and ionization potentials will therefore be estimated in a manner which is sufficiently general to cover these eventualities. A discussion of the results will be deferred to a following paper (Moffitt 1950).

2. SPECIFICATION OF HYBRID STATES

In this section a brief formulation of the various valence states in which we are interested will be given; it will also be convenient to introduce certain general descriptive terms in order to distinguish the different types of hybridization which occur.

TABLE 1

valence state	<i>A</i>	<i>B</i>	<i>C</i>	<i>D</i>	example
C, V_4	1	1	1	1	acetylene
C, V_2	2	1	1	0	carbon monoxide
N, V_3	2	1	1	1	cyanides
(N, $\bar{V}_3$)	1	1	2	1	isocyanides?)
O, V_2	2	1	2	1	carbon dioxide

The valence states with reference to fields of local symmetry $C_{\infty v}$ will first be specified. In the notation appropriate to this point group, the valence states envisaged are all of the form

$$(C_{\infty v}) V_n: (2u\sigma)^A (2t\sigma)^B (2p\pi_+)^C (2p\pi_-)^D, \quad (2.1)$$

where the sum of the occupation numbers ($A + B + C + D$) is 4, 5 or 6 according as we are dealing with carbon, nitrogen or oxygen respectively, and the number of valence electrons $n = \sum_T \delta_{T,1}$ ($T = A, B, C, D$). (Here, as in the remainder of this section, reference to the *K*-shell electrons will be omitted.) When the occupation number of a given orbital is 2, the corresponding electrons are formally non-bonding, and when it is equal to 1 the electron in this orbital is spoken of as a valence electron. The symbol V_n implies that the valence state V contains n valence electrons. The more interesting situations which arise are collected in table 1. The hybrid orbitals ($2u\sigma$) and ($2t\sigma$), which must be orthogonal to each other, are defined quite generally by the relations

$$\left. \begin{aligned} \phi(2u\sigma) &= \sin \chi \phi(2s) - \cos \chi \phi(2p\sigma), \\ \phi(2t\sigma) &= \cos \chi \phi(2s) + \sin \chi \phi(2p\sigma). \end{aligned} \right\} \quad (2.2)$$

Thus each of the valence states is uniquely specified by the set of occupation numbers (A, B, C, D) and a hybridization parameter χ . The energies of these different states may be calculated as functions of χ (see below) and, when referred to the ground states of their respective atoms, are called energies of promotion.

A typical valence state $V_n(\chi)$ may be said to arise from the atomic configuration $s^q p^r$, where

$$q = A \sin^2 \chi + B \cos^2 \chi, \quad r = A \cos^2 \chi + B \sin^2 \chi + C + D, \quad (2.3)$$

so that $(q+r) = (A+B+C+D)$. The configuration of the corresponding ground state is denoted by $s^{q'} p^{r'}$. (As for other many-electron systems, a density function ρ may be defined for $V_n(\chi)$; q and r are then easily shown to represent the average electron densities associated with $2s$ and $2p$ zones respectively.) The extension of this definition to more general valence states—of lower symmetry—is obvious.

When q, r are integers we speak of first-order hybridization in $V_n(\chi)$. If $q = q'$, the corresponding energy of promotion is small (case I), whereas if $q = q' - 1$ this energy term is much larger (case II) and corresponds mainly to the excitation $(2s)^{-1}(2p)$. The latter event is rarely encountered unless the valence number n of $V_n(\chi)$ exceeds the maximum valence number which can be attained in a configuration of the type $s^{q'} p^{r'}$. Examples of carbon valence states in classes I and II of first-order hybridization respectively are C, V_2 in CO , $^1\Sigma^+$ at large internuclear separations (Moffitt 1949) and C, V_4 in methane (Van Vleck 1934).

When q, r are non-integral, we speak of second-order hybridization in $V_n(\chi)$. Such cases only arise when the valence number n can be attained both by the configuration $s^{q'} p^{r'}$ and by $s^{q'-1} p^{r'+1}$. Their energies of promotion are intermediate between those for classes I and II of first-order hybridization. The conditions for second-order hybridization in divalent carbon, for example, have been examined by the author elsewhere (1949).

In the fields of lower symmetry C_{3v} and C_{2v} , appropriate valence states are easily written down:

$$(C_{3v}) V_n: (2ua_1)^A (2ta_1)^B (2pe_+)^C (2pe_-)^D, \quad (2.4)$$

$$(C_{2v}) V_n: (2ua_1)^A (2ta_1)^B (2pb_1)^C (2pb_2)^D, \quad (2.5)$$

using the standard notation for these point groups. The orbitals $(2ua_1), (2ta_1)$ are identical with $(2u\sigma), (2t\sigma)$ of (2.1) except that, for consistency of notation, $(2p\sigma)$ is now formally replaced by $(2pa_1)$ in (2.2). We shall prefer to treat fields of tetrahedral symmetry as particular examples of C_{2v} .

Very important cases occur when sets of orbitals $(2ta_1)(2pe_+)(2pe_-)$ of (2.4) and $(2ta_1)(2pb_1)$ of (2.5), which will be called $\mathbf{v}_3$ and $\mathbf{v}_2$ respectively, contain valence electrons. It is then easily shown that the members of $\mathbf{v}_m$ span an (m -dimensional) irreducible representation $\Gamma_m(\mathbf{v})$ of C_{mv} . It will often be found convenient, however, to replace $\mathbf{v}_3, \mathbf{v}_2$ by sets $\mathbf{h}_3, \mathbf{h}_2$ of mutually orthogonal, equivalent or directed orbitals. Members of $\mathbf{h}_m$ will span a reducible representation $\Gamma_m(\mathbf{h})$ of C_{mv} (cf. Van Vleck 1935). The group relation between $\mathbf{v}_m$ and $\mathbf{h}_m$ may therefore be expressed symbolically in the form

$$\mathbf{h}_m = S^{(m)} \mathbf{v}_m, \quad (2.6)$$

where $S^{(m)}$ may be determined as a unitary matrix which reduces $\Gamma_m(\mathbf{h})$ to $\Gamma_m(\mathbf{v})$. For C, V_4 in fields C_{2v} , $(2ua_1)$ and $(2pb_2)$ may also be hybridized in this way to obtain, finally, two pairs of equivalent orbitals with directional properties. More generally, the replacement (2.6) of any set of m valence orbitals by another set of m valence orbitals ($\mathbf{v}_m \rightarrow \mathbf{h}_m$), where $S^{(m)}$ may now be any m -dimensional unitary matrix, will be called 'local' hybridization. It is easily shown that local hybridization in any V_n does not alter the assignment of this V_n to a configuration $s^q p^r$.

Let us now see how the valence state $(C_{3v}) V_n$ is modified when $\mathbf{v}_3$ has been replaced by $\mathbf{h}_3$:

$$(C_{3v}) V'_n: (2ua_1)^4 (2h_1) (2h_2) (2h_3). \quad (2.7)$$

The matrix $S^{(3)}$ of (2.6) may be determined by means of group theory as

$$S^{(3)} = \begin{pmatrix} 1/\sqrt{3} & 1/\sqrt{3} & 1/\sqrt{3} \\ 1/\sqrt{3} & \frac{1}{2}(i - 1/\sqrt{3}) & \frac{1}{2}(-i - 1/\sqrt{3}) \\ 1/\sqrt{3} & \frac{1}{2}(-i - 1/\sqrt{3}) & \frac{1}{2}(i - 1/\sqrt{3}) \end{pmatrix}. \quad (2.8)$$

Explicit forms for the $(2h_i)$ are most simply referred to axes of quantization Oi ($i = 1, 2, 3$) taken along their directions of maximum electron density with the atomic nucleus as origin; $\phi(2p\sigma_i)$ implies that this orbital's component of angular momentum about Oi vanishes. We then have

$$\phi(2h_i) = \cos \varpi \phi(2s) + \sin \varpi \phi(2p\sigma_i) \quad (i = 1, 2, 3), \quad (2.9)$$

where the relation between χ and ϖ is most conveniently expressed in the form

$$3 \cos^2 \varpi = \cos^2 \chi. \quad (2.10)$$

It is clear from (2.10) that our construction of directed hybrid orbitals leaves invariant the quantities q, r of (2.3), so that our classification of the various types of hybridization which obtain remains valid. If we let $\beta = j\hat{O}i$ ($j \neq i$), we have

$$\cos \beta = -\cot^2 \varpi, \quad (2.11)$$

and, by use of the geometrical identity $\sin^2 \alpha = (\frac{4}{3}) \sin^2 (\frac{1}{2}\beta)$, we find that the angle α between Oi and the threefold symmetry axis is given by

$$\cos \alpha = \tan \chi \cot \varpi. \quad (2.12)$$

Similarly, we may replace $(2ta_1) (2pb_1)$ of $(C_{2v}) V_n$ by directed hybrids:

$$(C_{2v}) V'_n: (2ua_1)^4 (2h_1) (2h_2) (2pb_2)^D, \quad (2.13)$$

where the equations defining the $(2h_i)$ ($i = 1, 2$) are formally identical with (2.9). The corresponding matrix of transformation is

$$S^{(2)} = \begin{pmatrix} 1/\sqrt{2} & 1/\sqrt{2} \\ 1/\sqrt{2} & -1/\sqrt{2} \end{pmatrix}. \quad (2.14)$$

If α and β are the angles between Oi , the twofold axis of symmetry and Oj ($j \neq i$) respectively, these satisfy (2.12) and (2.11) once more. The hybridization parameters χ, ϖ are related by the equation

$$2 \cos^2 \varpi = \cos^2 \chi. \quad (2.15)$$

When $n = 4$, we may proceed further and replace $(2ua_1)(2pb_2)$ by the pair of equivalent orbitals

$$\phi(2h'_j) = \cos \lambda \phi(2s) + \sin \lambda \phi(2p\sigma) \quad (j = 3, 4), \quad (2.16)$$

to obtain the quadrivalent state

$$(C_{2v}) V_4'': (2h_1)(2h_2)(2h'_3)(2h'_4), \quad (2.17)$$

where

$$2 \cos^2 \lambda = \sin^2 \chi. \quad (2.18)$$

If, in addition, $\lambda = \varpi$ (i.e. $\chi = \frac{1}{4}\pi$), the hybridization is tetrahedral.

Examples of these different valence states will be given in a following paper (Moffitt 1950). Further states may also be constructed but are less interesting and reveal no new features of the hybridization process.

3. ENERGY RELATIONS

Let us determine the energy of the valence state

$$V_n: \dots (\phi_j)^J \dots \quad (j = 1, \dots, \nu), \quad (3.1)$$

where $J = N_j$ is the occupation number of ϕ_j , so that $0 < N_j \leq 2$, the total number of electrons $N = \sum_j N_j$ and $n = \sum_j \delta_{J,1}$. The ϕ_j are unspecified except in so far as they are all orthonormal $2s-2p$ hybrid orbitals (such as $\phi(2u\sigma)$ of (2.2)). Each singly occupied orbital, with $N_j = 1$, contains a valence electron the quantization of whose spin is independent of that of the remaining electrons. The eigenvalues of the scalar product operators $\mathbf{s}_j \cdot \mathbf{s}_i$ ($i \neq j$), representing its inclinations to the spins of the other electrons, all vanish.

By use of Dirac's vector formula, the energy of V_n may therefore be written in the form

$$E\{V_n\} = \sum_j N_j I_j + \sum_j (N_j - 1) J_{jj} + \sum_j \sum_{i>j} N_j N_i J_{ji} - \frac{1}{2} \sum_j \sum_{i>j} N_j N_i K_{ji} \quad (i, j = 1, \dots, \nu). \quad (3.2)$$

where, in atomic units,

$$I_j = \int \phi_j^* \left(-\frac{1}{2} \nabla^2 - \frac{Z}{r} \right) \phi_j dv + 2J_{jK} - K_{jK} \quad (3.3)$$

represents the energy of an electron of $(\phi_j)^J$ in the Hartree-Fock field of the nucleus, of charge Z , and the K -shell electrons $(\phi_K)^2$ (see appendix). We note that $(N_j - 1)$ can only assume the values 0 and 1. The Coulomb and exchange integrals are defined, as usual, by the relations

$$\left. \begin{aligned} J_{vw} &= \iint \phi_v^*(1) \phi_w^*(2) \frac{1}{r_{12}} \phi_v(1) \phi_w(2) dv_1 dv_2 \\ K_{vw} &= \iint \phi_v^*(1) \phi_w^*(2) \frac{1}{r_{12}} \phi_w(1) \phi_v(2) dv_1 dv_2 \end{aligned} \right\} \quad (v, w = K, 1, 2, \dots, \nu). \quad (3.4)$$

Thus when the hybridization parameters which define the ϕ_j have been laid down and the orbital functions $\phi(1s)$, $\phi(2s)$ and $\phi(2p)$ are known, the energy may be calculated using (3.2). As we shall later be using analytical approximations to these

orbital functions such calculations are not excessively laborious. Our results cannot be expected to be close to the accurate absolute values for the energies of binding $W\{V_n\}$ in these valence states, but we may reasonably expect calculations by means of (3.2) to enable us to estimate energies of promotion—and fortunately we may later determine to what extent this is so. At the moment, however, we content ourselves by noting that this absolute error $\Delta(\nu)$ will increase, numerically, with the number of electrons and that it is inherent in the use of our type of many-electron wave function.

Equation (3.2) may be rewritten in the form

$$E\{V_n\} = -\sum_j N_j \epsilon_j - \sum_j (N_j - 1) J_{jj} - \sum_j \sum_{i>j} N_j N_i J_{ji} + \frac{1}{2} \sum_j \sum_{i>j} N_j N_i K_{ji}, \quad (3.5)$$

$$\text{where} \quad -\epsilon_j = I_j + (N_j - 1) J_{jj} + \sum_i' N_i J_{ji} - \frac{1}{2} \sum_i' N_i K_{ji} \quad (i \neq j). \quad (3.6)$$

Comparison of (3.5) and (3.6) with (3.2) shows that $(E + \epsilon_j)$ has the same form as the energy $E\{V_n(\phi_j)^{-1}\}$, in our approximation, of the ion obtained by removing an electron in $(\phi_j)^J$ from V_n . (This is, of course, merely an extension of Koopmans's theorem to include valence states.) It will differ from the true energy of this ion in two respects. In the first place there will be the usual error $\Delta(\nu - 1)$ owing to the inadequacy of the general type of many-electron wave function employed. And in the second place our one-electron orbital functions will not be the best possible set for this ion, since their 'screening constants' $\{Z_i\}$ were determined by minimizing the energy of the neutral atom—incurring a second error $\Delta(Z^+)$. The ionization potential of ϕ_j may now be written

$$W\{V_n(\phi_j)^{-1}\} - W\{V_n\} = \epsilon_j + \{\Delta(\nu - 1) - \Delta(\nu)\} + \Delta(Z^+). \quad (3.7)$$

We note that whereas $\Delta(Z^+)$ is negative, $\{\Delta(\nu - 1) - \Delta(\nu)\}$ is positive. For atoms in spectroscopic states it has been found that the latter quantity is often nearly equal to $\Delta(Z^+)$, using the best (self-consistent) orbital functions. Accordingly, it is reasonable to expect that ϵ_j will be closer to the true ionization potential of ϕ_j than $E\{V_n(\phi_j)^{-1}\} - E\{V_n\}$, even though we use analytical approximations to the best orbital functions for V_n (see appendix). We shall therefore use equation (3.6), with certain modifications, in order to calculate the ionization potentials of hybrid orbitals. The validity of this procedure will be discussed more thoroughly later. Physically, $-\epsilon_j$ is an approximation to the energy of an electron of $(\phi_j)^J$ moving in the Hartree-Fock field of the remaining electrons and the nucleus.

Similarly the electron affinity of a valence orbital ϕ_j (with $N_j = 1$) may be written in the form

$$W\{V_n\} - W\{V_n(\phi_j)\} = \epsilon_j - J_{jj} + \{\Delta(\nu) - \Delta(\nu + 1)\} - \Delta(Z^-), \quad (3.8)$$

where $\Delta(Z^-)$, < 0 , represents the inadequate orbital representation in the negative ion $V_n(\phi_j)$. We see that $(\epsilon_j - J_{jj})$ is by no means as good an approximation to the electron affinity of ϕ_j as was ϵ_j to its ionization potential. Since $\{\Delta(\nu) - \Delta(\nu + 1)\}$ and $-\Delta(Z^-)$ are both positive, this quantity will be considerably less than the true value. Calculations of $(\epsilon_j - J_{jj})$ are nevertheless not without interest, as will be shown in Moffitt (1950).

Before proceeding, a lemma of some physical importance will be proved:

The value ϵ_j of the Hartree-Fock energy parameter associated with ϕ_j of V_n is independent of the particular state of local hybridization which may obtain among any set of valence orbitals

$$\{\phi_k\} (k \neq j) \text{ also of } V_n. \quad (3.9)$$

The proof is obvious. That part of ϵ_j which depends on the m valence orbitals comprising ϕ_k may be written in the form $\text{spur } \mathcal{M}^{(j)}$, where the elements of the m -dimensional matrix $\mathcal{M}^{(j)}$ are defined by

$$\mathcal{M}_{kl}^{(j)} = \iint \phi_j^*(1) \phi_k^*(2) \frac{1}{r_{12}} \phi_j(1) \phi_l(2) dv_1 dv_2 - \frac{1}{2} \iint \phi_j^*(1) \phi_k^*(2) \frac{1}{r_{12}} \phi_l(1) \phi_j(2) dv_1 dv_2. \quad (3.10)$$

The transformation $\phi_k \rightarrow \phi'_k = \sum_l S_{kl}^{(m)} \phi_l$, where $S^{(m)}$ is unitary, does not affect the value of $\text{spur } \mathcal{M}^{(j)}$, proving the lemma.

4. METHOD OF CALCULATION

The preceding section outlined the way in which energies of promotion and ionization potentials of hybrid valence states may be calculated in terms of certain integrals I_j , J_{ji} and K_{ji} . Since the hybrid orbitals ϕ_j are linear combinations of the simpler atomic orbitals $\phi(2lm_i)$, these integrals may be expressed in terms of the more primitive integrals $I(2l)$, $J(2lm_i; 2l'm'_i)$ and $K(2lm_i; 2l'm'_i)$, following the notation of Slater (1929). For example, when ϕ_j and ϕ_i are $\phi(2u\sigma)$ and $\phi(2t\sigma)$ of (2.2), we have

$$J(2u\sigma; 2t\sigma) = \sin^2 \chi \cos^2 \chi \{J(2s; 2s) + J(2p\sigma; 2p\sigma) - 4K(2s; 2p\sigma)\} + (\cos^4 \chi + \sin^4 \chi) J(2s; 2p\sigma). \quad (4.1)$$

In order to avoid the tedium of numerical integration, we shall use the best available analytical approximations to the Hartree-Fock atomic orbitals. (The results of Ufford's calculations (1938) indicate that the use of self-consistent orbital functions would not improve our calculations vastly.) These analytical functions have been determined by Duncanson & Coulson (1944). Their angular parts are the usual normalized tesseral harmonics and their radial parts have the form

$$\left. \begin{aligned} R(1s) &= L(1s) e^{-ar}, \\ R(2s) &= L(2s) (re^{-br} - r_0 e^{-dr}), \\ R(2p) &= L(2p) re^{-cr}. \end{aligned} \right\} \quad (4.2)$$

The $L(nl)$ are normalizing factors and r_0 is determined by the condition that $\phi(1s)$, $\phi(2s)$ be orthogonal. Values for the sets of remaining parameters a, b, c and d (which may, rather loosely, be called screening constants) have been determined for the atoms of the first period elements by minimizing the energies of their ground states. Strictly speaking we should, of course, use a different set of screening constants for each valence state we discuss—minimizing $E\{V_n\}$ with respect to these. This would complicate our calculations enormously, however, and it is doubtful whether such

a refinement is warranted by the greater errors inherent in the atomic orbital approximation. Accordingly, the sets of screening constants prescribed by Duncanson & Coulson will be used for all the valence states.

The 'one-electron' integrals appearing on the right-hand sides of equations (3.3) are now easily evaluated directly. Following Condon & Shortley's modification (1935) of Slater's procedure, we expand the $J(nlm_l; n'l'm_l')$ and $K(nlm_l; n'l'm_l')$ as follows:

$$\left. \begin{aligned} J(nlm_l; n'l'm_l') &= \sum_k a_k(lm_l; l'm_l') F_k(nl; n'l'), \\ K(nlm_l; n'l'm_l') &= \sum_k b_k(lm_l; l'm_l') G_k(nl; n'l'). \end{aligned} \right\} \quad (4.3)$$

Here the $a_k(lm_l; l'm_l')$ and $b_k(lm_l; l'm_l')$ are integers which have been tabulated as functions of their arguments (see, for example, Condon & Shortley 1935) and the F_k and G_k are defined by the relations

$$\left. \begin{aligned} F_k(nl; n'l') &= \frac{1}{D^k} \int_0^\infty \int_0^\infty R^2(nl/r_1) R^2(n'l'/r_2) r_a^k / r_b^{k+1} r_1^2 r_2^2 dr_1 dr_2, \\ G_k(nl; n'l') &= \frac{1}{D^k} \int_0^\infty \int_0^\infty R(nl/r_1) R(n'l'/r_1) R(nl/r_2) R(n'l'/r_2) (r_a^k / r_b^{k+1}) r_1^2 r_2^2 dr_1 dr_2, \end{aligned} \right\} \quad (4.4)$$

where r_a is the smaller, and r_b is the greater, of r_1 and r_2 . The D^k are certain denominators ensuring that the a_k, b_k be integral. This form of our integrals follows readily from Legendre's expansion of $1/r_{12}$ and the addition theorem for zonal harmonics. Since the orbitals of interest all have $l \leq 1$, the a_k, b_k vanish after $k = 2$. Using the sets of radial functions (4.2), the relevant integrations in (4.4) have therefore been performed for the carbon, nitrogen and oxygen atoms. The results are assembled in table 2.

TABLE 2. THE SLATER-CONDON PARAMETERS (IN ATOMIC UNITS)

parameter	carbon	nitrogen	oxygen
$-I(2s)$	2.27842	3.43508	4.82439
$-I(2p)$	1.92687	2.98335	4.23461
$F_0(2s; 2s)$	0.55192	0.65570	0.76057
$F_0(2s; 2p)$	0.57201	0.69084	0.80281
$F_0(2p; 2p)$	0.56672	0.69387	0.80648
$G_1(2s; 2p)$	0.12431	0.14958	0.17356
$G_2(2p; 2p)$	0.01097	0.01343	0.01561

To obtain some guide as to the validity of the procedure adopted, the calculated energies of the various states which arise from the configurations $s^a p^{r'}$ and $s^{a-1} p^{r'+1}$ of the atoms considered are tabulated, together with the observed energies of these states (table 3)—energies of all states being referred to their respective ground states. In this table, $G_1(2s; 2p)$ and $G_2(2p; 2p) = F_2(2p; 2p)$ are abbreviated to G_1 and F_2 respectively and W_0 is given by

$$W_0 = I(2p) - I(2s) + fF_0(2p; 2p) - F_0(2s; 2s) - gF_0(2s; 2p), \quad (4.5)$$

where for carbon, nitrogen and oxygen the sets of constants (f, g) take the sets of values (2, 1), (3, 2) and (4, 3) respectively. Tabulated energy quantities are give in electron-volts as distinct from the atomic units used in the text. The observed values are taken from Bacher & Goudsmit (1934).

TABLE 3. CALCULATED AND OBSERVED TERM VALUES IN C, N AND O
SPECTRA (IN UNITS eV)

atomic state	energy (eV)		
	theoretical	calculated	observed
C, s^2p^2 : 3P	—	—	—
1D	$6F_2$	1.79	1.26
1S	$15F_2$	4.48	2.68
sp^3 : 5S	$W_0 - G_1 - 10F_2$	3.46	4.16
3D	$W_0 - F_2$	9.52	7.94
3P	$W_0 + 5F_2$	11.31	9.33
1D	$W_0 + 2G_1 - F_2$	16.29	—
3S	$W_0 + 3G_1 - 10F_2$	16.98	13.12
1P	$W_0 + 2G_1 + 5F_2$	18.07	—
N, s^2p^3 : 4S	—	—	—
2D	$9F_2$	3.29	2.38
2P	$15F_2$	6.37	3.58
sp^4 : 4P	W_0	13.49	10.93
2D	$W_0 + G_1 + 6F_2$	19.75	—
2S	$W_0 + G_1 + 15F_2$	23.04	—
2P	$W_0 + 3G_1$	25.70	—
O, s^2p^4 : 3P	—	—	—
1D	$6F_2$	2.55	1.96
1S	$15F_2$	6.37	4.18
sp^5 : 3P	$W_0 + G_1 - 5F_2$	20.19	15.28
1P	$W_0 + 3G_1 - 5F_2$	25.38	—

The agreement is not startling but quite as good as can be expected, for the following reasons: In the first place, as has already been pointed out, the same set of orbital functions is used for the ground and the excited states of a given atom. Also, only first-order perturbation theory has been used—that is, in a terminology more suited to our approximation, inter-configurational interactions have been neglected. Thus the table shows that degrees of excitation are exaggerated by the calculations, in agreement with the sources of error cited.

It is interesting to note that the interval $s^2p^2, ^3P\ sp^3, ^5S$ for carbon is alone underestimated in the calculations—particularly since the set of screening constants were chosen so as to minimize the energy of a state arising from the configuration s^2p^2 . (The observed value given in table 3 is taken from Shenstone 1947.) This is not really surprising, since we should expect the orbital approximation to be the better suited as the multiplicity increases; the need for the introduction of correlation co-ordinates in a many-electron wave function should become less important as the multiplicity increases. The reality of this particular effect may be tested by comparing our interval (3.46 eV) with that obtained by use of self-consistent orbital functions, different sets being used for the s^2p^2 and sp^3 configurations respectively. The values calculated by the latter method should be even lower than the present figures. That

this is indeed the case is shown by Ufford's calculations of 1938, which give 3.16 eV. Ufford was unable to comment on this result, since the 5S state of carbon had not been observed at that date. We note, in this connexion, that on plotting the observed against the predicted degrees of excitation (as given in table 3) for carbon states of equal multiplicity, these states lie very nearly along straight lines. Such linearity cannot be shown for states of neutral nitrogen and oxygen, since an insufficient number of these have been observed.

TABLE 4. CALCULATED IONIZATION POTENTIALS AND MULLIKEN'S ESTIMATES (IN eV)

atomic state	ionic state	Mulliken	calculated ϵ_j
C (s^2p^2 , V_2)	C ⁺ (s^2p , V_1)	10.73	10.30
s^2p^2 , V_2	sp^2 , V_3	18.84	19.22
sp^3 , V_4	sp^2 , V_3	11.17	9.80
sp^3 , V_4	p^3 , V_3	(21.06)	20.37
N (s^2p^3 , V_3)	N ⁺ (s^2p^2 , V_2)	13.81	12.44
s^2p^3 , V_3	sp^3 , V_4	24.77	25.33
sp^4 , V_3	sp^3 , V_4	12.24	8.87
sp^4 , V_3	sp^3 , V_2	14.63	11.61
sp^4 , V_3	p^4 , V_2	(28.00)	26.41
O (s^2p^4 , V_2)	O ⁺ (s^2p^3 , V_3)	14.73	9.70
s^2p^4 , V_2	s^2p^3 , V_1	17.17	14.88
s^2p^4 , V_2	sp^4 , V_3	31.76	32.63

Similarly, in table 4 the values of ϵ_j are compared with the ionization potentials of atomic valence states—as estimated by Mulliken (1934*b*) from spectroscopic data. The spatial quantization of these valence states and of their ions is referred to an xyz system. The figures tabulated by Mulliken refer to the ionization potentials of pure p and s electrons from valence states in first-order hybridization and, apart from the rather less certain bracketed values, are correct to within 0.3 eV. Here again the agreement is not particularly good. We note that the ionization potential of an s electron is more nearly equal to the appropriate ϵ_j than is the case for a p electron. This is, of course, to be expected since the Hartree-Fock parameters are known to give much better estimates of the ionization potentials for inner shell than for valence electrons—and clearly s electrons are more strongly bound than p electrons.

From tables 3 and 4 we see, therefore, that equations (3.2) and (3.6) do not enable us to estimate energies of promotion and ionization potentials very accurately. We may, however, use the spectroscopic data to improve these estimates considerably. As Van Vleck (1934) has shown, energies of promotion for valence states in first-order hybridization may be calculated from observed spectra to within 0.3 eV. Since we have noted that the relation between observed degrees of excitation and corresponding predicted values, using equations analogous to (3.2), is approximately linear, the following procedure suggests itself:

Suppose we wish to calculate the energy of some $V_n(\chi)$ in second-order hybridization as a function of χ . Without loss of generality we may choose χ such that $q(\frac{1}{2}\pi) = q'$ and $q(0) = q' - 1$. Then we know (empirically) the values, $\bar{E}\{V_n(0)\}$ and

$\bar{E}\{V_n(\frac{1}{2}\pi)\}$, say, which this energy assumes for the values 0 and $\frac{1}{2}\pi$ of this argument. For intermediate values of χ , the formula

$$\bar{E}\{V_n(\chi)\} = \bar{E}\{V_n(\frac{1}{2}\pi)\} + \frac{\bar{E}\{V_n(0)\} - \bar{E}\{V_n(\frac{1}{2}\pi)\}}{E\{V_n(0)\} - E\{V_n(\frac{1}{2}\pi)\}} [E\{V_n(\chi)\} - E\{V_n(\frac{1}{2}\pi)\}], \quad (4.6)$$

will give a much better value for $W\{V_n(\chi)\}$ than does $E\{V_n(\chi)\}$ of (3.2). Our errors in predicting energies of promotion using (4.6) will not exceed 0.5 eV and are probably smaller than this. Intervals of the type $\bar{E}\{V_n(\chi_2)\} - E\{V_n(\chi_1)\}$, where $(\chi_2 - \chi_1)$ is small, will be very accurately determined by (4.6)—and it is in these that we shall be particularly interested.

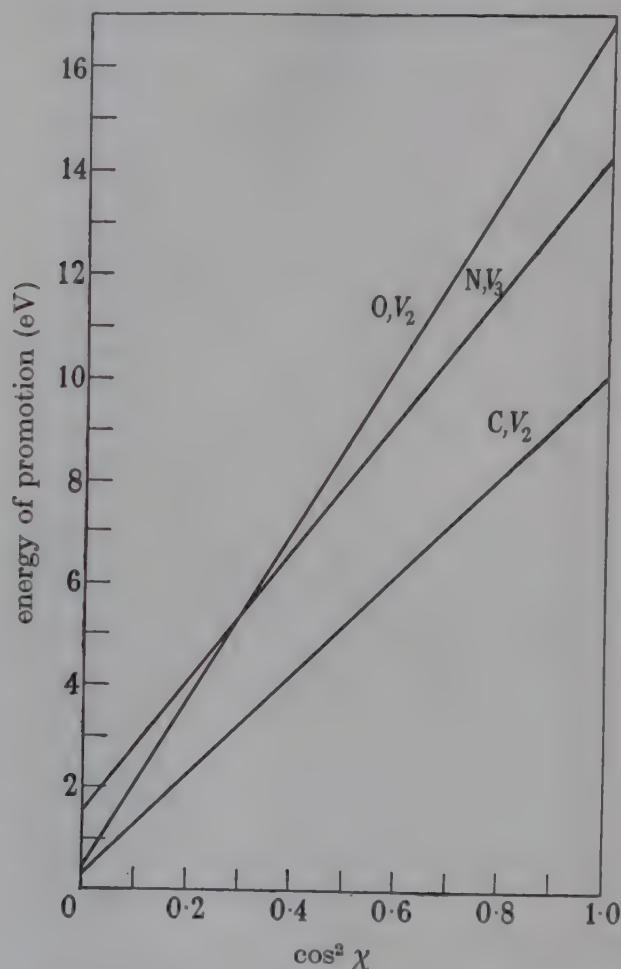


FIGURE 1. The energies of promotion for hybrid valence states.

Similarly, we may improve our estimates of the ionization potential of some hybrid orbital $\phi\{h(\chi)\} = \phi(2t\sigma)$ of (2.2), say, by means of the formula

$$\bar{\epsilon}\{h(\chi)\} = \bar{\epsilon}\{h(\frac{1}{2}\pi)\} + \frac{\bar{\epsilon}\{h(0)\} - \bar{\epsilon}\{h(\frac{1}{2}\pi)\}}{\epsilon\{h(0)\} - \epsilon\{h(\frac{1}{2}\pi)\}} [\epsilon\{h(\chi)\} - \epsilon\{h(\frac{1}{2}\pi)\}], \quad (4.7)$$

where $\bar{\epsilon}\{h(\frac{1}{2}\pi)\}$ and $\bar{\epsilon}\{h(0)\}$ are the 'observed' ionization potentials for pure p and s electrons respectively. The limits of accuracy for this formula are the same as those for (4.6).

In forming accurate estimates of both energies and ionization potentials of valence states, we have really used equations (3.2) and (3.6) as interpolation formulae. For some purposes, however, the ϵ_j and J_{jj} are useful in their own right. In these cases

they will be used to describe specific properties of Hartree-Fock fields which are, to that extent, independent of the validity of (3.2) and (3.6)—as will be shown in Moffitt (1950).

5. RESULTS

The main purpose of this paper has been to develop a method of calculating certain quantities which may be of some use in the analysis of molecular structure. Since the interpretation of our results is a matter for further analysis, a discussion of these is deferred to the following paper (Moffitt 1950). Only certain of the results

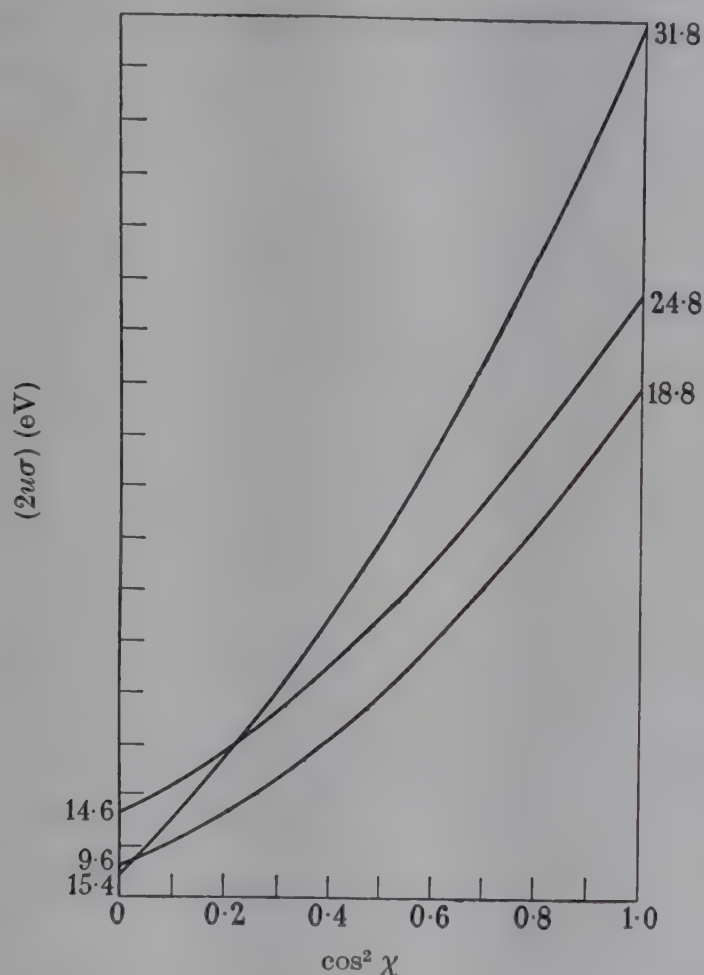


FIGURE 2. The ionization potentials of hybrid valence states. In order to bring all three curves on to the same graph, the origins for the ordinates are chosen differentially; the unit divisions represent eV.

will therefore be presented, and these graphically. The general tenor of the results is most easily seen by confining our attention to valence states in fields of local symmetry. $C_{\infty v}$. Figure 1 illustrates the way in which the energies of promotion (in eV) vary with the second-order hybridization parameter χ . These are computed by means of (3.2) and (4.6). The states considered are C, V_2 ; N, V_3 ; O, V_2 of table 1. Since the amount of $2s$ character of these states is linearly related to $\cos^2 \chi$ rather than to χ itself, the former quantities are used as abscissae. The curves are linear when plotted in this way.

Similarly, and for the same set of states, the ionization potentials $\bar{e}(2u\sigma)$ of the $(2u\sigma)$ electrons are plotted against $\cos^2 \chi$ in figure 2. The units are again eV. (Curves

may also, of course, be constructed for $\bar{\epsilon}(2t\sigma)$ —but these quantities are less directly related to observed properties of molecules and will not, for that reason, be reproduced here.

For other states, which may be derived from those of (2.1) by local hybridization, similar energy quantities may be estimated. In these cases the simplest procedure is to express integrals, such as $J(2h_i; 2h_j)$, linearly in terms of the integrals, like $J(2u\sigma; 2t\sigma)$ of (4.1), by means of the corresponding matrices $S^{(m)}$. Further, the results obtained may be extended to cover cases of partial ionicity. It is sometimes convenient to replace the integers B, C, D of (2.1) by fractional indices and, corresponding to such partially ionic valence states, we may calculate the ionization potentials of the non-bonding electrons ($2u\sigma$) by methods analogous to the use of (3.6) and (4.7).

APPENDIX

Formally the best atomic orbitals for any valence state V_n of the type (3.1) may be determined as the solutions of the integro-differential equations

$$\left\{ -\frac{1}{2}\nabla^2 - \frac{Z}{r} + 2 \int |\phi_K(2)|^2 \frac{1}{r_{12}} dv_2 + (N_j - 1) \int |\phi_j(2)|^2 \frac{1}{r_{12}} dv_2 + \sum_i' N_i \int |\phi_i(2)|^2 \frac{1}{r_{12}} dv_2 + \epsilon_{jj} \right\} \phi_j(1) - \int \phi_K^*(2) \frac{1}{r_{12}} \phi_j(2) dv_2 \phi_K(1) - \frac{1}{2} \sum_i' N_i \int \phi_i^*(2) \frac{1}{r_{12}} \phi_j(2) dv_2 \phi_i(1) + \sum_i' \epsilon_{ji} \phi_i(1) = 0, \\ (i, j = 1, \dots, \nu; i \neq j). \quad (\text{A } 1)$$

These equations may be derived by minimizing the energy $E\{V_n\}$ of (3.2) with respect to the one-electron orbital functions. From (3.2) the appropriate Euler variation equations are easily written down; the ϵ_{ji} appear as Lagrangian multipliers ensuring orthonormality and may be obtained in terms of the ϕ_j . For example, if we multiply (A 1) by ϕ_j^* and integrate throughout space, we find that

$$\epsilon_{jj} = -I_j - (N_j - 1)J_{jj} - \sum_i' N_i J_{ji} + \frac{1}{2} \sum_i' N_i K_{ji}. \quad (\text{A } 2)$$

Comparing this with (3.6) we see that ϵ_j and ϵ_{jj} are formally identical. In fact, of course, the ϵ_j which we have calculated are only approximations to the ϵ_{jj} , since we have used analytical approximations (4.2) to the self-consistent solutions of (A 1).

In conjunction with (A 1), we may define an operator $\mathcal{Z}\{V_n\}$, which is both linear and Hermitian, by means of the equation

$$(-\frac{1}{2}\nabla^2 + \mathcal{Z}\{V_n\} + \epsilon_{jj})\phi_j + \sum_i' \epsilon_{ji}\phi_i = 0 \quad (i, j = 1, \dots, \nu; i \neq j). \quad (\text{A } 3)$$

This operator $\mathcal{Z}\{V_n\}$ may be said to describe—for any electron in the M -shell—the Hartree-Fock field of the nucleus and the remaining electrons of V_n . (It is closely related to an operator first introduced by Dirac (1930) in a rather different context.) In this terminology, $-\epsilon_{jj}$ is the energy of an electron of $(\phi_j)^J$ moving in this field:

$$-\epsilon_{jj} = \int \phi_j^* (-\frac{1}{2}\nabla^2 + \mathcal{Z})\phi_j dv \quad (j = 1, \dots, \nu). \quad (\text{A } 4)$$

As a point of methodological interest, the derivation of the Hartree-Fock equations for atoms, molecules or metals is most simply and generally effected starting with Dirac's vector formula for the energy. This method is more direct than the group theoretical method on the one hand and avoids the explicit formulation of the many-electron wave function, which may involve several determinants, on the other. The advantages of our procedure are very apparent when applied to valence states and spectroscopic states, other than those corresponding to closed shell configurations and those of maximum multiplicity.

The work described in this paper forms part of a programme of fundamental research undertaken by the Board of the British Rubber Producers' Research Association. The author would like to thank Professor C. A. Coulson for reading the manuscript and Miss X. V. I. Sweeting for assistance with the numerical calculations.

REFERENCES

- Bacher, R. F. & Goudsmit, S. 1934 *Phys. Rev.* **46**, 948.
Condon, E. U. & Shortley, G. H. 1935 *Theory of atomic spectra*. Cambridge University Press.
Cottrell, T. L. & Sutton, L. E. 1947 *Quart. Rev. Chem. Soc.* **2**, 260.
Dirac, P. A. M. 1930 *Proc. Camb. Phil. Soc.* **26**, 376.
Duncanson, W. E. & Coulson, C. A. 1944 *Proc. Roy. Soc. Edinb.* **62**, 37.
Moffitt, W. 1949 *Proc. Roy. Soc. A*, **196**, 510, 524.
Moffitt, W. 1950 *Proc. Roy. Soc. A*, **202**, 548.
Mulliken, R. S. 1934a *Phys. Rev.* **46**, 551.
Mulliken, R. S. 1934b *J. Chem. Phys.* **2**, 782.
Pauling, L. 1931 *J. Amer. Chem. Soc.* **53**, 1367.
Shenstone, A. G. 1947 *Phys. Rev.* **72**, 411.
Slater, J. C. 1929 *Phys. Rev.* **34**, 1293.
Ufford, C. W. 1938 *Phys. Rev.* **53**, 568.
Van Vleck, J. H. 1934 *J. Chem. Phys.* **2**, 20, 297.
Van Vleck, J. H. 1935 *J. Chem. Phys.* **3**, 803.
Voge, H. H. 1936 *J. Chem. Phys.* **4**, 581.
Walsh, A. D. 1946 *Trans. Faraday Soc.* **42**, 56.

Aspects of hybridization

BY W. MOFFITT

British Rubber Producers' Research Association, Welwyn Garden City, Herts.

(Communicated by E. K. Rideal, F.R.S.—Received 14 March 1950)

It has been shown that the pairing theory of orbital hybridization accounts satisfactorily for the variations which are observed in the properties of C-H bonds. Since heteropolar effects are in the opposite directions to the effects described, it is concluded that hybridization influences the properties of predominantly covalent bonds to a greater extent than do differences in electronegativity. The extent of second-order hybridization in the molecules CH, NH and OH has been dealt with in the light of this analysis.

The consequences of the electronegativity concept have been examined on the basis of a generalized atomic orbital approximation. In particular, variations of the electronegativity of the carbon, nitrogen and oxygen atoms in different states of hybridization have been analyzed. The idea of 'lone' electrons has been formalized and the results of this definition have been discussed. Finally, the effects of orbital hybridization on the dipole moments and nuclear quadrupole coupling constants of molecules have been considered.

1. INTRODUCTION

In this paper it is intended to examine, rather generally, certain aspects of molecular structure and, in particular, those related to orbital hybridization. For simplicity, only those molecules for which localization is a good approximation will be treated. That is, it is assumed that only two electrons are responsible for the formation of a σ bond between two atoms. Inasmuch as the valence states of the constituent atoms are chosen for optimum bonding in the molecule as a whole, the use of the term 'localization' therefore only implies that hyperconjugation is neglected. Further, an atomic orbital approximation is used throughout, as is also implied in the use of the valence state concept. Many of the considerations developed below will be independent of the particular nature of this atomic orbital approximation. It is often immaterial whether the LCAO molecular orbital, or the perfect pairing AO method is used. In these cases, which include the consideration of ionic structures in the pairing approach on the one hand, and that of interconfigurational interactions in the molecular orbital method on the other, the approximation will be called that of generalized atomic orbitals (GAO). In other cases, when resort to a semi-empirical method must, however regrettably, be made, it will be found convenient to use the perfect pairing method. Although it has been verified that all the conclusions which are drawn in this way are in no way peculiar to this approximation, but could be reproduced in their general features by use of the molecular orbital method.

The nomenclature introduced in § 2 of the previous paper (Moffitt 1950—hereafter called I) will be used throughout.

2. THE STRENGTHS OF HYBRID BONDS

In this section it is intended to examine certain features of the pairing theory of carbon-hydrogen bonds, which has been used successfully by Van Vleck (1933), Penney (1935) and Coulson & Moffitt (1949). The aim of these authors has been

centred on the prediction of the shapes of molecules, rather than on the more intimate properties of individual C-H bonds. We shall, however, be more concerned with the latter aspects of the theory here. Our particular intention is to show that, although Pauling's criterion of bond strength cannot explain the observed variations of C-H bond strengths with changes in hybridization (Cottrell & Sutton 1947), a more thoroughgoing pairing treatment leads to results which are in complete qualitative agreement with the empirical results.

Consider the σ -bond formed by the pairing of the spins of a hybrid orbital

$$\phi_C(2t\sigma) = \cos \chi \phi_C(2s) + \sin \chi \phi_C(2p\sigma) \quad (2.1)$$

of quadrivalent carbon, and a $1s$ orbital, $\phi_H(1s)$, of hydrogen. Then the exchange energy associated with this bond may be written in the form

$$\begin{aligned} &\text{const.} + \left(\frac{3}{2}\right) K_{CH}(2t\sigma; 1s) \\ &= \text{const.} + \left(\frac{3}{2}\right) \iint \phi_C(2t\sigma | 1) \phi_H(1s | 2) \frac{1}{r_{12}} \phi_H(1s | 1) \phi_C(2t\sigma | 2) dv_1 dv_2, \end{aligned} \quad (2.2)$$

where the constant term is independent of the hybridization parameter χ , and of the state of local hybridization in the three remaining valence orbitals of carbon. In what follows, we shall relate C-H bond properties to this exchange term alone. The validity of our procedure has been discussed, for example, by Penney (1935).

The χ -dependent contribution of (2.2) to the C-H bond energy term may be rewritten, more enlighteningly, as

$$-\left(\frac{3}{2}\right) K_{CH}(2t\sigma; 1s) = \left(\frac{3}{2}\right) \{\sin^2 \chi N_{\sigma\sigma} + \sin 2\chi N_{s\sigma} + \cos^2 \chi N_{ss}\}, \quad (2.3)$$

where we have set

$$N_{\alpha\beta} = - \iint \phi_C(\alpha | 1) \phi_H(1s | 2) \frac{1}{r_{12}} \phi_H(1s | 1) \phi_C(\beta | 2) dv_1 dv_2 \quad (\alpha, \beta = 2s, 2p\sigma), \quad (2.4)$$

and, for convenience, abbreviated the subscripts $2s$ and $2p\sigma$ to s and σ respectively. Before proceeding, it is clearly necessary to discuss certain properties of the integrals $N_{\alpha\beta}$. These are, of course, functions of the internuclear separation $r(\text{C-H})$, and their magnitudes are fairly well established in the neighbourhood of $r(\text{C-H}) = 1.09 \text{ \AA}$. As a result of quadrature and of overlap considerations it may be shown that N_{ss} and $N_{\sigma\sigma}$ are both very close to 2 eV , as is $N_{s\sigma}$; Penney, for example, has given very good evidence for the set of values $N_{ss} = 2.0$, $N_{s\sigma} = 2.1$ and $N_{\sigma\sigma} = 2.2 \text{ eV}$. As a result of similar considerations, which are rather differently expressed by Coulson (1949) as differences in the covalent radii of the carbon $2s$ and $2p$ orbitals, we may make certain definite inferences about the derivatives of the $N_{\alpha\beta}$ with respect to the internuclear distance. A statement of these properties which is quite sufficient for our purposes may be given as follows:

$$\begin{aligned} N_{\alpha\beta} &\approx 2 \text{ eV} \quad (\alpha, \beta = s, \sigma), \\ \dot{N}_{ss} &\leq \dot{N}_{\sigma\sigma} < 0, \quad \dot{N}_{s\sigma} < 0, \end{aligned} \quad (2.5)$$

in the neighbourhood of $r(\text{C-H}) = 1.09 \text{ \AA}$, where the dots indicate differentiation with respect to $r(\text{C-H})$.

From the expression (2.3), it is now easily verified that the C-H bond energy term increases monotonically with $\cos^2 \chi$, the amount of $2s$ character in $\phi_C(2t\sigma)$, until it reaches its maximum value, when χ satisfies

$$\tan 2\chi = 2N_{s\sigma}/(N_{ss} - N_{\sigma\sigma}). \quad (2.6)$$

Substituting numerical values for the $N_{\alpha\beta}$, it is clear that the highest energy is attained by a bond for which χ lies very close to $\frac{1}{4}\pi$, or $\cos^2 \chi = \frac{1}{2}$.

As the carbon $2s$ character of a C-H bond increases from zero, so does the associated energy term until it reaches its maximum value in a bond formed with a hybrid carbon orbital of equal $2s$ and $2p$ character—after which it decreases. (2.7)

This result is very different from that of Pauling (1931), whose postulatory scheme would suggest that the strongest C-H bond is to be found in methane. It is much more in accord with Walsh's postulate (1948), that the C-H bond energy should increase with carbon $2s$ character over the whole range of hybridization—for C-H bonds with $\cos^2 \chi > \frac{1}{2}$ are unlikely to have more than a purely formal significance. Using Penney's set of values for the $N_{\alpha\beta}$, Förster obtained (2.7) rather differently in 1939. In the form (2.6), however, we see that (2.7) is particularly insensitive to an appropriate choice of values for the $N_{\alpha\beta}$. Further, it is easy to show that the equilibrium C-H internuclear separation decreases as $\cos^2 \chi$ increases, and also that the C-H stretching force constant increases with increasing carbon $2s$ character, providing the carbon atom valence state does not change greatly during vibration.

TABLE 1. THE PROPERTIES OF VARIOUS C-H BONDS

molecule	r_0 (C-H) (Å)	k (C-H)	
		(dynes/cm. $\times 10^6$)	D (C-H) (kcal.)
CH	1.131	4.3	80
CH ₄	1.094	5.0	102
C ₂ H ₄	1.071	5.1	100
C ₂ H ₂	1.059	5.9	104,121

In order to compare these results with observed properties of C-H bonds, we require an ordered series of molecules containing these bonds, for which $\cos^2 \chi$ may safely be assumed to increase monotonically. Such a series is furnished by the CH radical, methane, ethylene and acetylene respectively. For these molecules Walsh (1947*a*) has constructed a table, part of which is reproduced here (table 1). Walsh included a column giving the type of hybridization associated with each of these bonds. Since, however, there is some disagreement with his specification (see below), it has been omitted. The last column of table 1, it should be noticed, refers to bond dissociation energies, as distinct from bond-energy terms. Bond-energy terms are, of course, not thermodynamically well defined, and therefore not observable. The agreement between theory and experiment is entirely satisfactory. A more detailed, quantitative study of these factors will not be attempted—for such a programme would involve more than the study of localized C-H bonds—it being preferred to regard this qualitative agreement as evidence that the semi-empirical method, by means of

which it has been attained, is essentially correct. It is possible to account for the more intimate properties of bonds formed by hybrid orbitals, as well as the spatial configurations of the molecules in which they occur.

We may apply this method to the discussion of bending force constants. Penney (1935) has already considered a certain vibration (ν_2) of the methane molecule in this way, but it is more interesting to see how bending force constants vary with the hybridization parameter χ . To begin with the (unstrained) C-H bond with which we have been concerned above, suppose that the H atom is now moved in such a way that the C-H internuclear distance remains constant, but that its axis is inclined at an angle θ to the direction of maximum electron density of the carbon orbital $\phi(2t\sigma')$. ($\phi(2t\sigma')$ is written instead of $\phi(2t\sigma)$, because its $2p$ component, $\sin \chi \phi(2p\sigma')$, is now referred to an axis of quantization which does not coincide with the C-H direction, whereas the subscripts σ, π in the $N_{\alpha\beta}$ still refer to the C-H axis.) It is assumed that the carbon valence state does not 'follow' this motion to any extent, that is, that χ remains sensibly constant throughout the motion, and that interactions between non-nearest neighbours may be neglected. These conditions are not, of course, satisfied in actual molecules, but may be assumed to hold as a first approximation. Then the χ -dependent term in the potential energy function assumes the form

$$\left(\frac{3}{2}\right) K_{\text{CH}}(2t\sigma'; 1s) \\ = -\left(\frac{3}{2}\right) \{\cos^2 \chi N_{ss} + \sin 2\chi \cos \theta N_{s\sigma} + \sin^2 \chi \cos^2 \theta N_{\sigma\sigma} + \sin^2 \chi \sin^2 \theta N_{\pi\pi}\}; \quad (2.8)$$

$N_{\pi\pi}$ is defined by setting $\alpha = 2p\pi = \beta$ in (2.4)—it is known to have a value very close to -0.6 eV. The bending force constant $k_\theta(\chi)$ corresponding to this displacement is therefore the second derivative of the right-hand side of (2.8) with respect to θ , evaluated for the equilibrium position, which here corresponds to $\theta = 0$:

$$k_\theta(\chi) = \left(\frac{3}{2}\right) \{\sin 2\chi N_{s\sigma} + 2 \sin^2 \chi (N_{\sigma\sigma} - N_{\pi\pi})\}. \quad (2.9)$$

This expression attains its maximum value when

$$\tan 2\chi = -N_{s\sigma}/(N_{\sigma\sigma} - N_{\pi\pi}), \quad (2.10)$$

that is, when $\cos^2 \chi \approx 0.10$. For larger values of $\cos^2 \chi$, $k_\theta(\chi)$ decreases steadily and with increasing rapidity. Accordingly, we should expect values of $k_\theta(\chi)$ to be rather larger for bonds with $\cos^2 \chi = \frac{1}{4}$ (tetrahedral) than for those with $\cos^2 \chi = \frac{1}{3}$ (trigonal), and both to be appreciably greater than for those with $\cos^2 \chi = \frac{1}{2}$ (digonal). In the range of physical interest, therefore, we may say that angular deformations are more easily effected as the carbon $2s$ character of a C-H bond increases. This is, of course, easily explained, since the $2s$ orbital has no directional preferences. Although the exact interpretation of observed bending force constants requires a much more detailed analysis (cf. Linnett, Heath & Wheatley 1949), the figures quoted by these authors suggest that our simplified analysis is very pertinent. For the series $\text{C}-\text{C}/\text{H}$, $\text{C}=\text{C}/\text{H}$, $\text{C}\equiv\text{C}-\text{H}$ is one in which $\cos^2 \chi$ increases steadily from $\sim \frac{1}{4}$, and the bending constants are 0.66, 0.60 and 0.24 erg/radian², respectively.

There is another rather interesting conclusion which may be drawn in this way. Suppose a carbon atom forms a σ -bond with some other atom, say nitrogen. And

further, let this σ -bond be compressed to an internuclear separation which is appreciably shorter than the equilibrium single-bond distance. Then the criterion for strongest C-N bonding is no longer the analogue of (2.7), but is shifted in favour of greater carbon $2s$ character. This is easily demonstrated since the derivatives of the corresponding primitive exchange integrals will satisfy inequalities which are formally analogous to those in (2.5). With an extended definition of the $N_{\alpha\beta}$, the condition (2.6) for the best bond still holds. By differentiating with respect to the internuclear separation r , this becomes

$$\frac{\partial \chi}{\partial r} = - \left\{ \frac{N_{s\sigma}(\dot{N}_{ss} - \dot{N}_{\sigma\sigma}) - \dot{N}_{s\sigma}(N_{ss} - N_{\sigma\sigma})}{4N_{s\sigma}^2 + (N_{ss} - N_{\sigma\sigma})^2} \right\}. \quad (2.11)$$

And, in view of (2.5), we have

$$\frac{\partial}{\partial r} (\cos^2 \chi) = -\sin 2\chi \frac{\partial \chi}{\partial r} \approx \sin 2\chi \frac{(\dot{N}_{ss} - \dot{N}_{\sigma\sigma})}{4N_{s\sigma}} < 0, \quad (2.12)$$

when χ lies in the interval $(0, \frac{1}{2}\pi)$, proving our corollary. Our result has certain interesting applications: in the first place, it suggests that the C-H bonds in acetylene contain slightly less $2s$ character than is generally supposed; for the (compressed) σ -component of the $\text{C}\equiv\text{C}$ bond will tend to increase its share of this. Secondly, second-order hybridization at the carbon atom will be more important in the carbon monoxide molecule (Moffitt 1949), than in the CH radical (see below). Finally, it provides a formal justification for some of the qualitative arguments which Coulson, Duchesne & Manneback (1948) have advanced in their discussion of the signs of interaction constants.

3. SECOND-ORDER HYBRIDIZATION

It is generally assumed that the bonding in, for example, the CH radical ($^2\Pi$) is of pure $2p$ type as regards the carbon valence orbital. On the other hand, it is also recognized that $2s - 2p$ hybridization increases the C-H bond strength, but that such hybridization involves an energy of promotion at the carbon atom. It is of obvious interest, therefore, to investigate the truth of the original assumption. Such calculations have not been carried out hitherto, because estimates of second-order hybridization energies for atoms such as carbon have not been available. The analysis in I now makes it possible for us to determine the extent of second-order hybridization in the molecules CH, NH and OH. The calculations are simple, since it is only necessary to determine where the opposing demands for stable intra-atomic and inter-atomic bonding lead to an optimum molecular energy. In order to assess variations in the inter-atomic energy terms, the pairing method is used because, as we have seen in the preceding section, it appears that this method is particularly useful for explaining the way in which hybridization modifies these terms.

The appropriate C, N and O valence states are the V_n of (I, 2.1), namely,

$$(C_{\infty v}) V_n(\chi): (2u\sigma)^2 (2t\sigma) (2p\pi_+)^C (2p\pi_-)^D, \quad (3.1)$$

where $(C + D) = 1, 2$ and 3 for these atoms respectively. The valence state energies are given as functions of the hybridization parameter χ by the relation

$$\text{const.} + P_M \cos^2 \chi \quad (M = \text{C, N, O}), \quad (3.2)$$

where $P_M = 9.75, 12.90$ and 16.40 eV respectively (cf. figure 1 of I). When $\cos^2 \chi = 0$, the M -H bonding is pure $2p$. The validity of (3.2) is unaffected by the particular way in which the spin-orbital degeneracy of the $2p\pi$ electrons is resolved in the case of the nitrogen atom.

The χ -dependent inter-atomic exchange term is in all cases of the form

$$K_{MH}(2t\sigma; 1s) - K_{MH}(2u\sigma; 1s) \quad (M = \text{C, N, O}), \quad (3.3)$$

where the K_{MH} are the usual exchange integrals. More explicitly, (3.3) may be written as

$$-\{\cos 2\chi (N_{ss}^M - N_{\sigma\sigma}^M) + 2 \sin 2\chi N_{s\sigma}^M\}. \quad (3.4)$$

We have extended our definition of the primitive M -H exchange integrals to include atoms other than carbon by an obvious generalization of (2.4). Coulomb terms are again assumed to be independent of χ —the justification for which is rather less well founded than in the case of quadrivalent carbon. It should not lead us into serious error, however.

Treating χ as a parameter with respect to which the total molecular energy must now be minimized, we have

$$\frac{\partial}{\partial \chi} \{P_M \cos^2 \chi - \cos 2\chi (N_{ss}^M - N_{\sigma\sigma}^M) - 2 \sin 2\chi N_{s\sigma}^M\} = 0. \quad (3.5)$$

The optimum conditions may therefore be obtained in the simple form

$$\tan 2\chi = -4N_{ss}^M / \{P_M - 2(N_{ss}^M - N_{\sigma\sigma}^M)\} \quad (M = \text{C, N, O}). \quad (3.6)$$

The amount of $2s$ character in the M -H bond is simply $\cos^2 \chi$, where χ is now that root of (3.6) which lies in the interval $(0, \frac{1}{2}\pi)$. In all cases, $2(N_{ss}^M - N_{\sigma\sigma}^M)$ may be neglected in comparison with P_M and $N_{s\sigma}^M$. Accordingly the additional energy of stabilization due to second-order hybridization is

$$2 \sin 2\chi N_{s\sigma}^M - P_M \cos^2 \chi. \quad (3.7)$$

Exact values of the $N_{s\sigma}^M$ are unfortunately not available; it is, however, known that they will be not less than 1.5 eV and probably not appreciably greater than 2.0 eV. Accordingly $\cos^2 \chi$ and the quantity (3.7) have been evaluated in each case, taking the possible values $1.50, 1.75$ and 2.00 eV for the $N_{s\sigma}^M$. The results are assembled in table 2. We note that second-order hybridization is quite significant in all cases, but more so for carbon than for oxygen. In particular, the CH radical contains about 10 % of carbon $2s$ character, which contributes some 30 kcal. to the total energy of formation.

TABLE 2. SECOND-ORDER HYBRIDIZATION IN MH

molecule	$N_{s\sigma}^M$ (eV)	amount of $2s$ character	gain in energy (eV)
CH	1.50	0.0741	0.849
	1.75	0.0939	1.126
	2.00	0.1132	1.434
NH	1.50	0.0467	0.664
	1.75	0.0605	0.889
	2.00	0.0751	1.140
OH	1.50	0.0305	0.532
	1.75	0.0402	0.716
	2.00	0.0506	0.924

4. ELECTRONEGATIVITIES

In the preceding sections it has been tacitly assumed that hybridization affects the strengths of σ -bonds (in particular C-H, N-H and O-H) to a greater extent than do departures from homopolar character in these bonds. That such departures do occur is abundantly clear. It seems equally clear that these play a subsidiary role in determining the strengths of predominantly covalent σ -bonds, such as occur in saturated hydrocarbons. It is of interest, however, to see how hybridization affects the heteropolar character of bonds, and it is to this question that the present section is primarily devoted.

We shall not lose in generality by considering a σ -bond τ , between atoms a and b of a polyatomic molecule, in the pairing symbolism—considering structures a^+b^- , $a-b$ and a^-b^+ . For, in the last resort, a molecular orbital description involving interconfigurational interactions is equivalent to this. The wave function representing τ may be written in the form

$$\Psi^\tau = \gamma_1^\tau \Psi(a^+b^-) + \gamma_2^\tau \Psi(a-b) + \gamma_3^\tau \Psi(a^-b^+). \quad (4.1)$$

When the γ^τ have been determined, we may associate resultant formal charges with atoms a and b . If $\gamma_1^\tau > \gamma_3^\tau$, the formal (negative) charge associated with b is greater than that associated with a , and we say that b is more electronegative than a under these circumstances. (The phases of the structure functions are chosen in such a way that the Ψ are all real and the γ^τ are all positive.)

More precisely, let ϕ_a^τ , ϕ_b^τ , be the τ -bonding orbitals of atoms a and b respectively; in general these will be hybrids. Then the formal charge separation in τ is associated with the unequal partition of two electrons between ϕ_a^τ and ϕ_b^τ . To this extent, differences in hybridization will affect the electronegativity exhibited by an atom under different conditions. In general, therefore, we shall need to specify the valence states of atoms a and b before we can say which of these is more electronegative with regard to the bond τ . As an example, for a given value of χ , a carbon atom in the state $(C_{\infty v})V_2(\chi)$ will show a different electronegativity with respect to hydrogen, than it will in the state $(C_{\infty v})V_4(\chi)$. And similarly its electronegative properties in $(C_{\infty r})V_4(\chi)$, say, will depend on χ .

Clearly, therefore, it is not sufficient to attach a single number, representing its electronegativity, to an atom irrespective of its environment. It is, however, in this sense that the term electronegativity has most generally been applied. We shall suppose that it is sufficient to specify the appropriate atomic valence states in order to discuss their electronegative propensities. This specification is most certainly necessary, as we have already seen. Its sufficiency derives from that of the use of the term 'electronegativity'. For if the electronegativity of an atom is more intimately related to its environment, much of the utility of this concept disappears.

Suppose, therefore, that the valence states of a and b have been determined by symmetry or otherwise (e.g. by the higher valence considerations discussed in §3 above). We assume that the distribution of those electrons of the molecule which are not involved in the bond τ , is known; for simplicity, we shall assume it to be homopolar. The Hamiltonian determining Ψ^τ may be written in the form

$$\mathcal{H} = \mathcal{H}_0 + \mathcal{Z}_a^\tau(1) + \mathcal{Z}_b^\tau(1) + \mathcal{Z}_a^\tau(2) + \mathcal{Z}_b^\tau(2) + 1/r_{12} - \frac{1}{2}\nabla^2(1) - \frac{1}{2}\nabla^2(2); \quad (4.2)$$

for convenience, we have called the two τ -bonding electrons 1 and 2, without jeopardizing the fundamental Pauli antisymmetry. $\mathcal{L}_a^\tau(1)$, for example, represents the Hartree-Fock field of the nucleus and the non- τ -bonding electrons of atom a , and is defined for operation on $\phi_a^\tau(1)$ and $\phi_b^\tau(1)$ (cf. appendix in paper I). Terms in the Hamiltonian, such as the interactions of the non- τ -bonding electrons of atom a with those of b , which do not affect the subsequent argument, have been subsumed under $\mathcal{H}_0$. The best *GAO* representation of the bond τ is now obtained by minimizing the energy $\int \Psi^\tau \mathcal{H} \Psi^\tau dv$, subject to the normality condition, $\int \Psi^\tau \Psi^\tau dv = 1$. We shall use the same atomic orbitals in forming Ψ^τ as were used in discussing atomic energy levels in I—that is, we assume that it is sufficient here to use undistorted atomic orbitals.

The electronegativity concept is now introduced in the following form: Orbital ϕ_a^τ is said to be more or less electronegative than orbital ϕ_b^τ according as $\gamma_1^\tau < \text{or} > \gamma_3^\tau$ respectively. Further, it is assumed that this qualitative criterion of electronegativity is dependent only on properties of the (pre-determined) valence states of atoms a and b —and is therefore independent of the $a-b$ internuclear separation. (This does not, of course, commit us to the more radical and unjustified assumption that the formal charges associated with atoms a and b do not depend on $r(a-b)$.) Our assumption is regarded as fundamental to the electronegativity concept and, so far as the author is aware, is not contravened by any experiment. It may be tested, at least, in principle, for diatomic molecules (in which complicating factors such as π -bonding effects are absent), by plotting their dipole moments $\mu(a-b)$ as functions of their internuclear separations $r(a-b)$; the heteropolar contribution to $\mu(a-b)$ should then only vanish at infinity or constantly throughout the open interval

$$r(a-b) \geq r_e(a-b).$$

Let us therefore determine the criterion for equal electronegativities of ϕ_a^τ and ϕ_b^τ under these conditions. By minimizing the energy, it is found that $\gamma_1^\tau = \gamma_3^\tau$ for all values of $r(a-b)$ if, and only if,

$$\begin{aligned} & \int \phi_a^{\tau*}(1) \left\{ -\frac{1}{2} \nabla^2(1) + \mathcal{L}_a^\tau(1) \right\} \phi_a^\tau(1) dv_1 + \frac{1}{2} \iint |\phi_a^\tau(1)|^2 \frac{1}{r_{12}} |\phi_a^\tau(2)|^2 dv_1 dv_2 \\ &= \int \phi_b^{\tau*}(1) \left\{ -\frac{1}{2} \nabla^2(1) + \mathcal{L}_b^\tau(1) \right\} \phi_b^\tau(1) dv_1 + \frac{1}{2} \iint |\phi_b^\tau(1)|^2 \frac{1}{r_{12}} |\phi_b^\tau(2)|^2 dv_1 dv_2, \quad (4.3) \end{aligned}$$

always provided that the electronegativity assumption is valid. The demonstration is elementary but a little tedious, so that it will be omitted. In the notation of I, this may be written in the more enlightening form

$$\epsilon(\phi_a^\tau) - \frac{1}{2} J(\phi_a^\tau; \phi_a^\tau) = \epsilon(\phi_b^\tau) - \frac{1}{2} J(\phi_b^\tau; \phi_b^\tau), \quad (4.4)$$

where $\epsilon(\phi_a^\tau)$ is the Hartree-Fock energy parameter associated with ϕ_a^τ in the appropriate valence state of atom a and $J(\phi_a^\tau; \phi_a^\tau)$ is the usual Coulomb integral. Further, it may be shown that if the quantity on the right-hand side of (4.4) is less than that on the left-hand side, then $\gamma_1^\tau > \gamma_3^\tau$, and vice versa. Clearly, therefore, the quantities $\epsilon(\phi_a^\tau) - \frac{1}{2} J(\phi_a^\tau; \phi_a^\tau)$ have all the properties with which we wish to characterize absolute

electronegativities. Their only limitations in this context derive from those of the valence state approximation, which presupposes the GAO approximation which we have used.

A relation which is analogous to (4.4) was derived by Mulliken (1949) under less general conditions using the molecular orbital method. Our derivation is more closely related to his more qualitative considerations of 1934. Apart from the more general lines of the above analysis, it is hoped that we have been able to throw a clearer light on the electronegativity concept as a whole. The relation of this approach to that of Pauling (1939) is shelved for the moment; Pauling's electronegativity scale is based on his interpretation of the as yet unsubstantiated postulate of the additivity of covalent bond energies. Recent work by Sutton & Cottrell (1950) throws some doubt on its validity and, in any case, Pauling's scale is not nearly sensitive enough for our purposes. More profitably, we may contrast the different ways in which the relation (4.4) may be interpreted.

Mulliken observed that $\epsilon(\phi_a^\tau) - \frac{1}{2}J(\phi_a^\tau; \phi_a^\tau)$ is related to the average of the ionization potential and electron affinity $\frac{1}{2}(I + E)$, of the orbital ϕ_a^τ . For in a certain approximation, which we have discussed in §3 of I, $\epsilon(\phi_a^\tau)$ may be identified with this ionization potential and $\epsilon(\phi_a^\tau) - J(\phi_a^\tau; \phi_a^\tau)$ with the electron affinity. Accordingly he regards (4.4) as a more rigorous foundation for the highly successful electronegativity scale which he set up in 1934. In this connexion, however, it should be remarked that $\epsilon(\phi_a^\tau)$ is not a very good approximation to the 'observed' ionization potential, and we have given good reasons for supposing that

$$\epsilon(\phi_a^\tau) - J(\phi_a^\tau; \phi_a^\tau)$$

is a much poorer approximation to the electron affinity (cf. I, 3.8).

Another interpretation of (4.4) remains open to us: we may regard (4.4) as furnishing us with the definition of a purely theoretical electronegativity scale, based on the valence state approximation. For, as we have seen in I, values of the quantities appearing in this equation are amenable to calculation. Such a procedure should certainly be a reliable guide for assessing the way in which the electronegativity of an atom depends on its environment (state of hybridization). It is in this sense that the relation (4.4) will be used. Before applying it to the comparison of the electronegativities of two different atoms, however, it is necessary to be satisfied that the valence state approximation is equally valid in both cases. When this information is lacking, it is probably safer to use the $\epsilon(\phi_a^\tau) - \frac{1}{2}J(\phi_a^\tau; \phi_a^\tau)$ as a means of interpolation in conjunction with Mulliken's proved scale (1934). For the latter requires interpolation formulae when hybrid orbitals are involved.

With this reservation, we may now calculate the theoretical electronegativities of different hybrid valence orbitals as functions of their 2s character, specifying the valence states to which they refer. The abbreviation

$$\eta(\phi_a^\tau) = \epsilon(\phi_a^\tau) - \frac{1}{2}J(\phi_a^\tau; \phi_a^\tau) \quad (4.5)$$

will be used and values of the η given in atomic energy units. For carbon, the values of $\eta(2t\sigma)$ for the valence states V_2 and V_4 , respectively, referred to fields of local symmetry $C_{\infty v}$ have been calculated (cf. I, 2.1). The results are illustrated in figure 1,

where the abscissae are values of $\cos^2 \chi$, the amount of $2s$ character in $\phi(2t\sigma)$. The electronegativity is a monotonically increasing function of $\cos^2 \chi$ in both cases, and is linear for V_4 . The lemma (I, 3.10) shows that the values of $\eta(2t\sigma)$ for V_4 apply whatever the condition of local hybridization may be in the three remaining valence orbitals of carbon. That is, the lower curve in figure 1 applies to all hybrid orbitals of quadrivalent carbon. We notice that, for a given amount of $2s$ character, divalent carbon is more electronegative than quadrivalent carbon. For nitrogen, the values $\eta(2t\sigma)$'s have been calculated for $(C_{\infty v}) V_3(\chi)$ of (I, 2.1) and $\eta(2h_i)$'s for $(C_{3v}) V'_3(\chi)$ of (I, 2.7). The abscissae again denote amounts of $2s$ character which are given, respectively, by $\cos^2 \chi$ and $\frac{1}{3} \cos^2 \chi$ for these orbitals. The results are represented graphically in figure 2. Similarly, figure 3 shows $\eta(2t\sigma)$ and $\eta(2h_i)$ for the respective states $(C_{\infty v}) V_2(\chi)$ and $(C_{2v}) V'_2(\chi)$, of (I, 2.13), of the oxygen atom. The amounts of $2s$ character in the valence orbitals of interest are $\cos^2 \chi$ and $(\frac{1}{2}) \cos^2 \chi$ respectively.

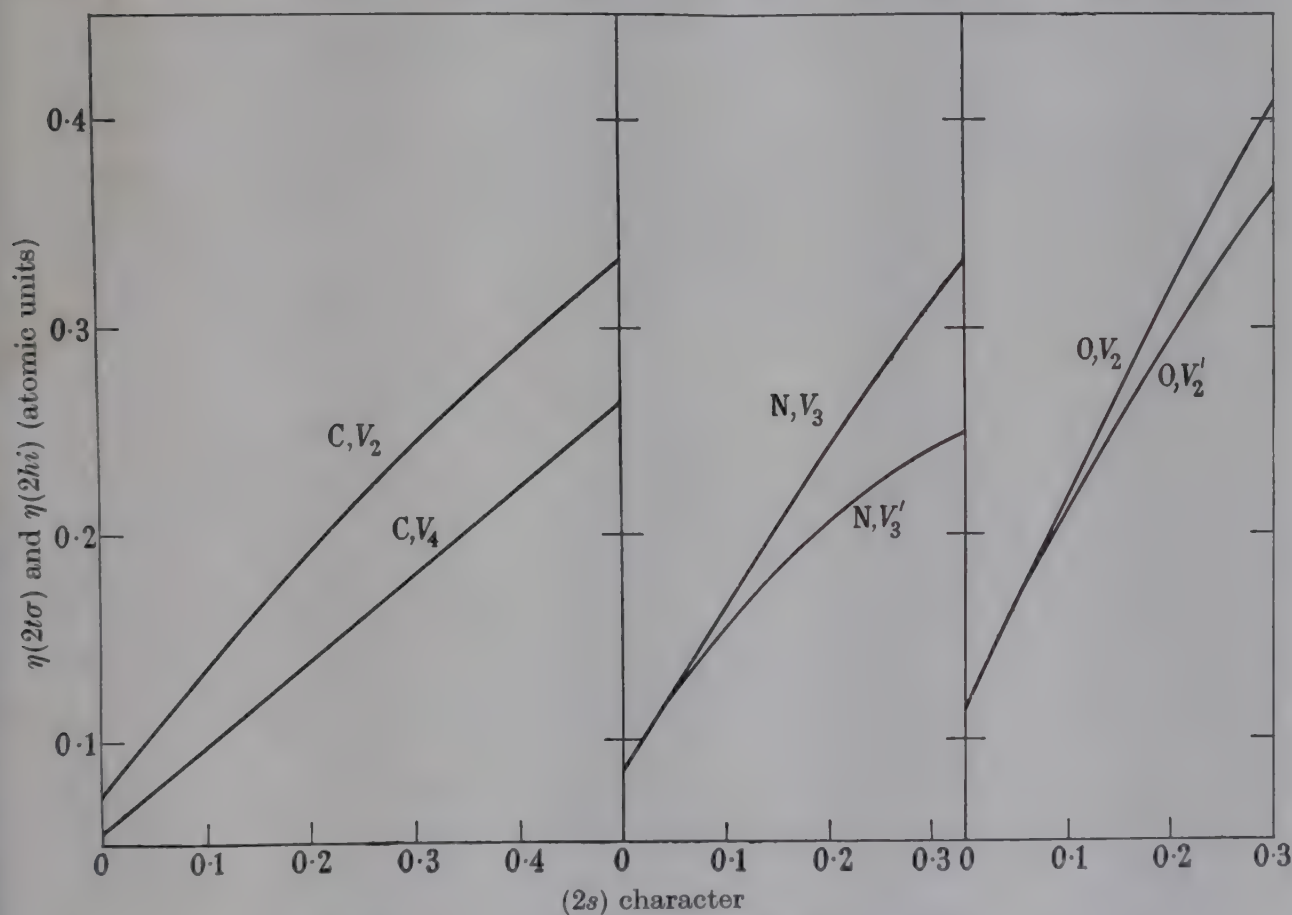


FIGURE 1

FIGURE 2

FIGURE 3

FIGURE 1. The electronegativities of hybrid carbon orbitals.

FIGURE 2. The electronegativities of hybrid nitrogen orbitals.

FIGURE 3. The electronegativities of hybrid oxygen orbitals.

It is of interest that, on this scale, the electronegativity of the hydrogen atom $\eta(1s)$ is 0.1875 atomic units. Second-order hybridization here—e.g. the inclusion of hydrogen $2p$ character to strengthen bonding (Rosen 1931)—will serve to decrease this value somewhat. For whereas second-order hybridization in the C, N and O atoms is, to a first approximation, associated with the partial inclusion of the more stable $2s$ orbitals in bonding (and therefore with an increase in electronegativity), the

acquisition by $\phi_H(1s)$ of increased directional properties involves the participation of less stable orbitals such as $\phi_H(2p\sigma)$. Hence the electronegativity exhibited by the hydrogen atom will be ~ 0.186 atomic units.

We have, therefore, a purely theoretical guide as to the way in which the electronegativity of an atom depends on its state of hybridization. Further work remains to be done before we may compare the electronegativities of a given carbon orbital with that of a given oxygen orbital, say, on such theoretical grounds. For, as we have already indicated, a much more complete knowledge of the way in which the limitations of the valence state concept vary, on going across the periodic table, is necessary before this can be attempted. Again, the author believes that the formal charges which are at present associated with atoms in even the simplest of molecules are probably erroneous (see below), so that further analysis of the empirical data is also required before a suitable comparison between theory and experiment may be made. We hope to discuss this topic more fully at a later date.

The variations in electronegativity which we have found, and estimated quantitatively, were first discussed by Mulliken in 1934. More recently, Walsh (1947*a*) has applied similar qualitative arguments to correlate a large body of experimental material. To the principle that the electronegativity of a hybrid orbital increases with its $2s$ character, he has also added the following corollary: 'If a group X attached to carbon is replaced by a more electronegative group Y , then the carbon valency towards Y has more $2p$ character than it had towards X .' This subsidiary principle, according to Walsh, is a translation into quantum-mechanical language of the well-known charge transfer or inductive effect. It is easily seen, however, that this is by no means the case. For the inductive effect of polarities in a bond τ formed by carbon is reflected in another bond ρ , made by the same C atom, by the dependence of $\mathcal{Z}_C$ on the charge density in ϕ_C^e , at least to a first approximation. Analysis shows that the factor embodied in Walsh's corollary is a tendency for the carbon atom to fill its $2s$ zone before its $2p$ zone. However, unless the heteropolar bonding in C- Y becomes comparable with the covalent component of the total C- Y bonding, it is easily shown that this effect is only a very secondary factor in determining the appropriate carbon valence state. The recent work of Gordy, Simmons & Smith (1948) completely substantiates the results of Penney's calculations (1935), who used only the covalent bonding terms, on the geometry of the methyl halides. We conclude that, except in cases of very high polarity, the best-bonding factors, considered in §2 above, are more important than those arising from applications of Walsh's corollary.

5. LONE SETS OF ELECTRONS

As is well known, it often happens that certain pairs of molecular electrons may be assigned to particular atoms of that molecule. This assignment may be made, either on grounds of symmetry (e.g. the (π_g) electrons of CO_2 , $^1\Sigma_g^+$) or on grounds of plausibility (e.g. the $(u\sigma)$ electrons of CO , $X^1\Sigma^+$). In our GAO approximation, this concept may be formalized as follows:

If in any molecule the orbital motion of a certain set of electrons may, to good approximation, be represented by an orbital of a particular atom, and the orientation

of the vector representing the total spin of these electrons is independent of the orientations of the remaining spin vectors, then this set of electrons is said to constitute a 'lone' set on that atom.

Very frequently these sets of electrons occur in pairs, when we speak of 'lone pairs' of electrons. The orbital to which the n^μ electrons of a lone set μ are assigned will, in general, be a hybrid orbital ϕ_c^μ of the atom c with which they are associated.

We may now investigate some of the implications of this concept of 'lone' electrons. We shall not lose in generality by concentrating on a lone pair $(\phi_c^\mu)^2$, where ϕ_c^μ is real and non-degenerate—the extension to other cases being obvious. Since these electrons are adequately represented by an atomic orbital of c , we must assume that the remaining electrons which are associated with c in the molecule effectively screen the lone pair in such a way that the first-order energy of the molecule M may be written in the simple form

$$E\{M\} = E\{M(\sim\mu)\} + 2 \int \phi_c^\mu \left\{ -\frac{1}{2}\nabla^2 + \mathcal{Z}_c^\mu \right\} \phi_c^\mu dv + \iint |\phi_c^\mu(1)|^2 \frac{1}{r_{12}} |\phi_c^\mu(2)|^2 dv_1 dv_2, \quad (5.1)$$

where the constant term $E\{M(\sim\mu)\}$ in $E\{M\}$ is independent of ϕ_c^μ , and $\mathcal{Z}_c^\mu$ represents the Hartree-Fock field of the c atom nucleus and the remaining electrons on c . For, only if the Coulomb-exchange field of the electrons in M which are associated with nuclei other than c , and of these nuclei themselves, effectively vanish throughout the lone-pair region, will an unperturbed hybrid orbital ϕ_c^μ of c be at all apposite.

The Hartree-Fock energy parameter associated with an electron of $(\phi_c^\mu)^2$ is

$$-\epsilon(\phi_c^\mu) = \int \phi_c^\mu \left\{ -\frac{1}{2}\nabla^2 + \mathcal{Z}_c^\mu \right\} \phi_c^\mu dv + \iint |\phi_c^\mu(1)|^2 \frac{1}{r_{12}} |\phi_c^\mu(2)|^2 dv_1 dv_2, \quad (5.2)$$

so that $E\{M\} + \epsilon(\phi_c^\mu)$ has the same form as the energy $E\{M(\phi_c^\mu)^{-1}\}$ of the ion obtained by removing one of the lone electrons from M (in our approximation). $E\{M\} + \epsilon(\phi_c^\mu)$ will differ from the true energy $W\{M(\phi_c^\mu)^{-1}\}$ of this ion in what we may conveniently regard as three different respects. We may once more distinguish between an error $\Delta(\nu-1)$ due to the general type of wave function employed and $\Delta(Z^+)$ due to an inappropriate choice of screening constants for the atomic orbitals in their changed (ionic) environment (cf. I, 3.7). But in dealing with molecules we encounter another error which will be called $\Delta(\gamma^+)$ and which arises in the following manner: The term $E\{M(\sim\mu)\}$ in (5.1) contains a number of parameters $\{\gamma\}$ in addition to the screening constants $\{Z\}$. These parameters are chosen so as to minimize the energy of M in some approximation (cf. §4). Now $\{\gamma\}$, which incidentally is also a set of the arguments determining $\mathcal{Z}_c^\mu$, will not be an optimum set for $M(\phi_c^\mu)^{-1}$ and this gives rise to the third error. The appropriate ionization potential is

$$W\{M(\phi_c^\mu)^{-1}\} - W\{M\} = \epsilon(\phi_c^\mu) + \{\Delta(\nu-1) - \Delta(\nu)\} + \Delta(Z^+) + \Delta(\gamma^+). \quad (5.3)$$

Granted the validity of the lone pair assignment, we may now make certain inferences concerning this ionization potential. With as much justification as for atoms we may either calculate $\epsilon(\phi_c^\mu)$ and neglect $\{\Delta(\nu-1) - \Delta(\nu) + \Delta(Z^+)\}$, or, more reliably, estimate the sum of these terms, which are called $\bar{\epsilon}(\phi_c^\mu)$ —as explained in I.

We then have

$$W\{M(\phi_c^\mu)^{-1}\} - W\{M\} = \bar{\epsilon}(\phi_c^\mu) + \Delta(\gamma^+). \quad (5.4)$$

Since $\Delta(\gamma^+)$ is, by definition, negative we see that $\bar{\epsilon}(\phi_c^\mu)$ gives us an upper limit to the ionization potential of an electron ϕ_c^μ of the molecule M . (5.5)

The magnitude of this inequality, represented by $\Delta(\gamma^+)$, provides us with a measure of the reorganization of the molecule on ionization. If we use the so-called 'vertical' potentials, in which the internuclear separations remain essentially unaltered on ionization, this energy of reorganization is almost entirely electronic in origin.

In using (5.5) we must remember that the operators $\mathcal{Z}_c^\mu$ will, in general, refer to fields associated with partially ionic, or non-integral valence states—as distinct from the neutral valence states, whose occupation numbers are integers, with which we were largely concerned in I. This arises because of the possibility that the bonds, made by c with adjacent atoms, are to some extent heteropolar. An appropriate $\bar{\epsilon}(\phi_c^\mu)$ may always be estimated when the electron distribution around c is known.

Our result (5.5) suggests that Walsh's estimates (1947*b*) of formal charge distributions in a series of ketones by means of ionization potential data are not quite correct. He has tacitly identified $\bar{\epsilon}(\phi_c^\mu)$ with observed ionization potentials, and consequently exaggerated the formal charge effects—always provided that his lone pair assignment is justified. Again, observed ionization potentials less than the corresponding $\bar{\epsilon}(\phi_c^\mu)$'s have been interpreted as indications of anti-bonding characteristics in the ϕ_c^μ . Deduction (5.5) shows that this is not necessarily the case.

It would, of course, be very interesting to know how nearly the condition (5.1) is fulfilled in actual molecules. Unfortunately, this requires a complete knowledge of the molecular charge distributions, and these are not available. We may, however, see whether (5.5) is satisfied by typical non-bonding electrons, and whether the corresponding energies of reorganization appear reasonable. In order to simplify our identification, we choose two molecules in which the 'lone' electrons are known to be formally non-bonding on grounds of symmetry. The methyl radical may be shown, on very convincing theoretical grounds, to be planar. It contains an electron, (a_2''), in the notation appropriate to the point group D_{3h} , which is non-bonding. Its formal charge distribution is homopolar to good approximation. On these grounds, $\bar{\epsilon}(a_2'') \approx 11$ eV. Hipple & Stevenson (1943) have observed a vertical ionization potential of 10.0 eV. The corresponding energy of reorganization, ~ 1 eV, is quite consistent with the sort of rigidity we would associate with the CH_3 radical. Similarly, Mulliken (1935) has shown that the four (π_g) electrons of carbon dioxide are best regarded as localized in non-bonding pairs, with one pair on each of the two oxygen atoms. An approximate molecular orbital method (Moffitt 1949) suggests that the charge distribution in this molecule is nearly homopolar. Accordingly $\bar{\epsilon}(\pi_g) \approx 15$ eV. Since the observed ionization potential is 13.7 eV, we again find an energy of reorganization in the neighbourhood of 1 eV. (We note that Mulliken's estimate of a formal charge of -0.54 on each of the oxygen atoms is inconsistent with the assignment of the (π_g) electrons to lone sets, in our sense. For in this case we should have $\bar{\epsilon}(\pi_g) \approx 9$ eV, which contravenes (5.5).)

These examples, however plausibly interpreted, illustrate the difficulties which are inevitably encountered in attempting to analyze observed ionization potentials. In the absence of any more complete knowledge of molecular charge distributions, it seems best to make the heuristic assumption that non-bonding electrons may be

regarded as lone, in this specialized sense, and to use the evidence of (5.5) in conjunction with data derived from other sources. In a related series of molecules, uncertainties of interpretation are of course, minimized.

6. PHYSICAL MANIFESTATIONS OF HYBRIDIZATION

Second-order hybridization may manifest itself physically in several different ways, two of which we shall discuss here. Both of these may be considered as electrical in origin and demonstrate the asymmetry of the electronic charge distribution in certain valence states. It is easily shown that they are unaffected by changes in local hybridization, so that we may, with little loss of generality, confine ourselves to fields of symmetry $C_{\infty v}$.

The dipole moments of molecules may be conveniently regarded as made up from three contributions. Two of these do not concern us here (see, however, §4), namely, the heteropolar moments due to partial ionicity and the moments which are induced by these in other regions of the molecule. An important part of the third, homopolar contribution to the measured dipole moment arises from the so-called atomic dipoles characteristic of valence states in second-order hybridization. In general a valence state of the type (I, 3.1) will be associated with a dipole moment owing to its intrinsic asymmetry in charge distribution. Taking the direction of the initial line in our system of polar co-ordinates as positive, this dipole moment will be of magnitude

$$\mu\{V_n\} = \sum_j N_j \int \phi_j^* r \cos \theta \phi_j dv \quad (6.1)$$

in atomic units. For $V_n(\chi)$ of (I, 2.1) this becomes

$$\mu\{V_n(\chi)\} = -\sin 2\chi(A - B) \int \phi(2s) r \cos \theta \phi(2p\sigma) dv. \quad (6.2)$$

The atomic dipole of such a $V_n(\chi)$ therefore only vanishes in first-order hybridization—that is, when $\chi = 0$ or $\frac{1}{2}\pi$, or $A = B$.

This term is of obvious importance to the theoretical interpretation of the dipole moments of molecules such as NH_3 and H_2O —it was first estimated quantitatively in the carbon monoxide molecule (Moffitt 1949), where it was shown to be very significant. The integral which appears in (6.2) is accordingly evaluated:

In the notation of I,

$$\int \phi(2s) r \cos \theta \phi(2p\sigma) dv = \frac{1}{\sqrt{3}} L(2s) L(2p) \int_0^\infty R(2s) r R(2p) r^2 dr. \quad (6.3)$$

The latter integral is easily obtained so that, in Debye units,

$$e \int \phi(2s) r \cos \theta \phi(2p\sigma) dv = 2.31, 1.91 \text{ and } 1.65 \quad (6.4)$$

for the carbon, nitrogen and oxygen atoms respectively. (In a previous paper (1949), the author gave the value 2.4 Debye units for the carbon integral on the basis of less accurate atomic functions.) For partially ionic states, the screening constants in the $R(nl)$ will need modification. For a given value of χ , the numerical value of integral

(6.3) will decrease slightly as the positive charge associated with the atom increases. In view of the approximations involved in our GAO model, it is however doubtful if such variations will be significant.

Owing to recent developments in the technique of micro-wave spectroscopy, it has become possible to investigate a new set of molecular properties, the so-called nuclear quadrupole coupling constants. These are properties of molecules containing a nucleus whose departure from spherical symmetry, as gauged by its electric quadrupole moment Q , gives rise to an interaction between Q and the electrostatic field V of the remaining molecular charges. Since the possible orientations of such a nucleus on the principal (z -) axis of a linear or symmetric top molecule are quantized with respect to this axis, it has been possible to measure nuclear quadrupole coupling constants, defined by the relation

$$eQq = eQ(\partial^2 V / \partial z^2), \quad (6.5)$$

for these molecules. Townes & Dailey (1949) have shown that in many cases, some 90 % or more of the term q in (6.5) arises from the electrons which, in a given molecule, are associated with the atomic nucleus whose quadrupole moment is in question—and that this eventuality often provides us with a sensitive guide as to the properties of atomic valence states. At the moment we shall merely be interested in assessing the contribution to q of these electrons, and to give formulae for the variations of this contribution with changes of hybridization and polarity.

Let us consider an atomic valence state V_n of the type (I, 3.1), whose nucleus has a non-vanishing quadrupole moment Q . We shall be thinking of the N^{14} nucleus in particular. Then the contribution of the electrons of V_n to the nuclear quadrupole coupling constant is

$$eQq\{V_n\} = eQ \sum_j N_j \int \phi_j^* \frac{3 \cos^2 \theta - 1}{r^3} \phi_j dv = eQ \sum_j N_j q_{jj}, \quad \text{say,} \quad (6.6)$$

where we have used the relation

$$\frac{\partial^2}{\partial z^2} \left(\frac{1}{r} \right) = \frac{3z^2 - r^2}{r^5} = \frac{3 \cos^2 \theta - 1}{r^3}.$$

(The K -shell electrons may, to a very good approximation, be neglected.) When V_n is not partially ionic, the N_j are integers and it is easily shown (cf. I, 3.10) that the expression on the right-hand side of (6.6) is invariant with respect to changes of local hybridization. Under such circumstances, therefore, no information as to the extent of delocalization effects at a particular V_n can be obtained in this way. Hyperconjugation cannot be detected in the absence of polarities—and only very indirectly when these are present (vide infra). Thus Townes & Dailey's tentative interpretation of a series of eQq measurements as evidence in favour of the delocalized description

$N \equiv H_3$ of the ammonia molecule, as against the localized $H-N \begin{smallmatrix} \nearrow H \\ \searrow H \end{smallmatrix}$, cannot be considered as valid.

Since we are largely concerned with linear or symmetric top molecules, we may therefore specialize our discussion to the $V_n(\chi)$ of (I, 2.1), which are defined for fields of symmetry $C_{\infty v}$. In this case (6.6) becomes

$$q\{V_n(\chi)\} = Aq(2u\sigma; 2u\sigma) + Bq(2t\sigma; 2t\sigma) + (C + D)q(2p\pi_+; 2p\pi_+), \quad (6.7)$$

where we have replaced q_{jj} by $q(2u\sigma; 2u\sigma)$ when $\phi_j = \phi(2u\sigma)$ and so on. In view of certain simple identities which obtain among the q_{ji} , (6.7) becomes

$$q\{V_n(\chi)\} = (A \cos^2 \chi + B \sin^2 \chi) q(2p\sigma; 2p\sigma) + (C + D) q(2p\pi; 2p\pi). \quad (6.8)$$

Further, utilizing the symmetry properties of the angular parts of the $2p$ orbital functions, we have

$$q(2p\sigma; 2p\sigma) = 2q_0(2p; 2p) \quad \text{and} \quad q(2p\pi; 2p\pi) = -q_0(2p; 2p),$$

where

$$q_0(2p; 2p) = \left(\frac{2}{5}\right) \int_0^\infty R^2(2p) r^{-3} r^2 dr = \left(\frac{2}{5}\right) \langle r^{-3} \rangle. \quad (6.9)$$

(Our notation is an obvious extension of that used in I for the discussion of energy properties.) (6.8) now reads

$$\begin{aligned} eQq\{V_n(\chi)\} &= eQq_0(2p; 2p) \{2(A \cos^2 \chi + B \sin^2 \chi) - (C + D)\} \\ &= \kappa \{2(A \cos^2 \chi + B \sin^2 \chi) - (C + D)\}, \end{aligned} \quad (6.10)$$

where κ is a constant. This formula applies to neutral and to partially ionic valence states alike. (In cases of large polarity, however, the factor $q_0(2p; 2p)$ in κ will be affected to some extent, owing to the changes in the screening constant associated with $R(2p)$. A correction for this factor is easily applied whenever necessary.) Variations of the quadrupole coupling constant measured for a non-spherical nucleus in a variety of molecular environments will therefore be associated, very largely, with changes in the non-constant term on the right hand side of (6.10).

Let us now consider a nitrogen nucleus in a field of symmetry C_{3v} . We replace $V_3 \equiv (C_{3v}) V_3(\chi)$ of (I, 2.4) by $V'_3 \equiv (C_{3v}) V'_3(\chi)$ of (I, 2.7) by local hybridization, and compare V_3 and V'_3 in partial ionicity. For V_3 , there is now no necessary relation between non-integral B and $(C + D)$. On the other hand, if electron pairing is a good approximation, V'_3 should be appropriate and the relation $B = \frac{1}{2}(C + D)$ should hold. When χ , $eQq_0(2p; 2p)$ and the total formal charge associated with the nitrogen atom are known, departures from $B = \frac{1}{2}(C + D)$ are theoretically observable by means of nuclear quadrupole coupling constants and would be a measure of the degree of delocalization. Since, however, it is most unlikely that such small effects could unambiguously be assigned to valence state effects, that χ , eQq_0 and the formal charge are at all precisely determined, and that such a simplified atomic orbital approximation is capable of this sort of precision, it is quite safe to say that delocalization effects cannot be detected in this way.

7. CONCLUSIONS

The arguments which have been followed in the preceding sections may be summarized as follows:

(i) The electron pairing theory of orbital hybridization has been shown to account, very satisfactorily, for the variations which have been observed in the properties of C-H bonds. From this it may be concluded that hybridization is a more important factor in determining the strengths of such predominantly covalent bonds than is their partially ionic character, for example. Further, we have seen that second-order hybridization in the comparatively simple molecules CH, NH and OH is quite

significant; in each of these, it accounts for an additional stabilization which is of the order of 1 eV (§§ 2, 3).

(ii) An investigation of the electronegativities of atoms in hybrid states has been undertaken. In the course of this, we have been led to regard the electronegativity concept from a rather different perspective to that which is customary (§ 4).

(iii) When electrons are assigned to so-called non-bonding orbitals of particular atoms, certain properties are generally attributed to these. In the belief that current usage is to some extent contradictory, a concept of 'lone' electrons on particular atoms has been formalized. The consequences of this definition have been discussed. If sets of non-bonding electrons in hybrid orbitals satisfy the condition that they be 'lone', in our sense, then their ionization potentials, taken in conjunction with certain calculated quantities, will provide useful information about the structures of the molecules in which they occur (§ 5).

(iv) Finally, the effect of hybridization on the dipole moments and nuclear quadrupole coupling constants of molecules have been discussed. The results show, *inter alia*, that the crude assignment of formal charges in any molecule, on the basis of dipole moment and bond distance measurements alone, is very unreliable when atoms of that molecule are capable of second-order hybridization (§ 6).

It is hoped that these considerations may assist in clarifying certain puzzling features of molecular properties to which Cottrell & Sutton drew attention in their recent review article (1947).

The work described in this paper forms part of a programme of fundamental research undertaken by the Board of the British Rubber Producers' Research Association. The author would like to thank Professor C. A. Coulson for reading the manuscript and Miss X. V. I. Sweeting for assistance with the numerical calculations.

REFERENCES

- Cottrell, T. L. & Sutton, L. E. 1947 *Quart. Rev. Chem. Soc.* **11**, 260.
 Cottrell, T. L. & Sutton, L. E. 1950 To be published.
 Coulson, C. A., Duchesne, J. & Manneback, C. 1948 *Victor Henri Memorial Volume*. Liège: Desoer.
 Coulson, C. A. & Moffitt, W. 1949 *Phil. Mag.* **40**, 1.
 Förster, T. 1939 *Z. phys. Chem. B*, **43**, 58.
 Gordy, W., Simmons, J. W. & Smith, A. G. 1947 *Phys. Rev.* **72**, 344.
 Hipple, J. A. & Stevenson, D. P. 1943 *Phys. Rev.* **63**, 121.
 Linnett, J. W., Heath, D. F. & Wheatley, P. J. 1949 *Trans. Faraday Soc.* **45**, 832.
 Moffitt, W. 1949 *Proc. Roy. Soc. A*, **196**, 510, 524.
 Moffitt, W. 1950 *Proc. Roy. Soc. A*, **202**, 534.
 Mulliken, R. S. 1934 *J. Chem. Phys.* **2**, 782.
 Mulliken, R. S. 1935 *J. Chem. Phys.* **3**, 720.
 Mulliken, R. S. 1949 *J. Chim. phys.* **46**, 497.
 Pauling, L. 1931 *J. Amer. Chem. Soc.* **53**, 1367.
 Pauling, L. 1939 *The nature of the chemical bond*. New York: Cornell.
 Penney, W. G. 1935 *Trans. Faraday Soc.* **31**, 734.
 Rosen, N. 1931 *Phys. Rev.* **38**, 2099.
 Townes, C. H. & Dailey, B. P. 1949 *J. Chem. Phys.* **17**, 782.
 Van Vleck, J. H. 1933 *J. Chem. Phys.* **1**, 177, 219.
 Walsh, A. D. 1947a *Faraday, Soc. Disc.* 'The labile molecule', p. 18.
 Walsh, A. D. 1947b *Trans. Faraday Soc.* **43**, 158.

Surf beats: sea waves of 1 to 5 min. period

BY M. J. TUCKER, *Royal Naval Scientific Service*

(Communicated by G. E. R. Deacon, F.R.S.—Received 25 March 1950)

Long sea waves with periods of 2 to 3 min. and a few inches in amplitude have been measured. It has been shown that they are due to the varying height of groups of waves breaking on the shore. The amplitude of the long waves is found to be approximately proportional to the amplitude of the ordinary waves and independent of their period. The mechanism by which they are produced is discussed.

INTRODUCTION

At the conference of the International Union of Geodesy and Geophysics in Oslo in August 1948, W. H. Munk reported that he had detected fluctuations in sea level of a few inches' amplitude with periods of 2 to 20 min. He attributed the waves of about 2 min. period to variations in the height of the surf and called them 'beats of surf'. He considered that the waves of longer period, about 20 min., might come directly from distant storms, and since waves of this length would travel at several hundred miles an hour, they might be used for the detection and location of storms.

Munk (1948) used a pneumatic apparatus, the maximum sensitivity of which was to waves of 15 min. period with a pass band extending from 2 min. to 2 hr.; it was designed to be insensitive to ordinary sea waves and to the tides, even though their amplitudes are much greater than those of the 2 min. to 2 hr. waves. The apparatus consisted of a series of vessels partly filled with oil and connected by capillary tubes. For the measurements described in this paper it was more convenient to use a wave-pressure recorder of the 'aneroid' type already laid approximately 1000 yards from the beach at Perranporth on the north coast of Cornwall. The voltage from this instrument (about 0.5 mV/ft.) was amplified by a simple photoelectric d.c. amplifier with negative feedback. The output was passed through two resistance-condenser couplings of very long time-constant, designed to remove the tidal components, into a 2 min. period galvanometer slightly less than critically damped and insensitive to ordinary sea-wave frequencies. The circuit arrangement is shown in figure 1, and the relative response of the system to pressure fluctuations of 12 sec. to 20 min. period in figure 2. Owing to a leaky air valve in the undersea unit, the overall response decreased rapidly for periods longer than 5 min. Records of long waves, using the long-period galvanometer, and of ordinary waves, using a galvanometer of $\frac{1}{2}$ sec. period, were made side by side on photographic paper $5\frac{1}{2}$ in. wide moving at $\frac{1}{4}$ in./min. Time marks were recorded every 20 min.

Unfortunately, the 'aneroid' wave recorder—the only one at Perranporth suitable for measuring pressure fluctuations longer than 30 sec. period—was becoming rapidly unserviceable owing to the very slow leakage of air from the rubber bag,

and it was possible to obtain satisfactory records only for about 1 hr. on either side of low water. Records were obtained from 6 to 28 December, after which the unit, which had been working for 18 months, failed entirely. The records are, however,

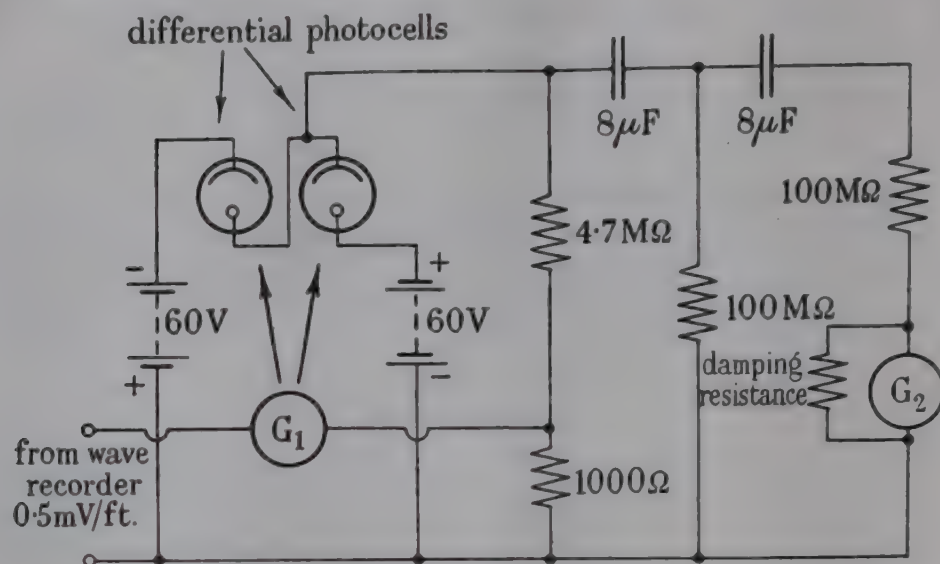


FIGURE 1. Long-wave filter circuit, incorporating simple negative-feedback d.c. amplifier, anti tidal-drift coupling and long-period galvanometer insensitive to ordinary-wave periods. G_1 is a sensitive galvanometer and G_2 a 2 min. period galvanometer.

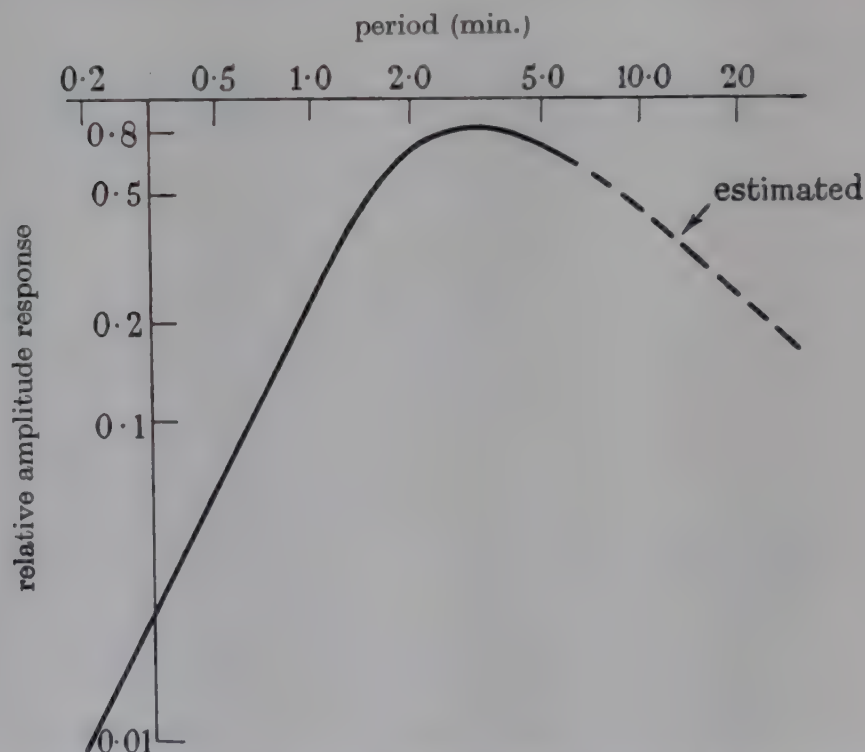


FIGURE 2. Frequency response of the recorder system, taking into account the leaky air valve in the undersea unit.

adequate for investigation of waves of about 3 min. period. No waves of longer period were detected, but the decreased sensitivity of the apparatus for waves longer than 5 to 10 min. period leaves room for doubt as to whether they were present.

RELATIONSHIP BETWEEN LONG AND ORDINARY WAVES

On 23 December 1948 a calm period was interrupted by the arrival of swell from a storm which developed off Newfoundland (2000 miles from Perranporth) on 21 December. The peak heights of the long and ordinary waves are plotted against time in figure 3. The peak-wave height is the difference in height between the crest and trough of the highest wave on the length of record examined. It is an easy measurement to make and experience has shown that it is approximately proportional to the average wave height. There is an approximately linear relationship between the amplitudes of the long waves and ordinary waves. This is further illustrated in figure 4 by plotting the two measurements against one another.

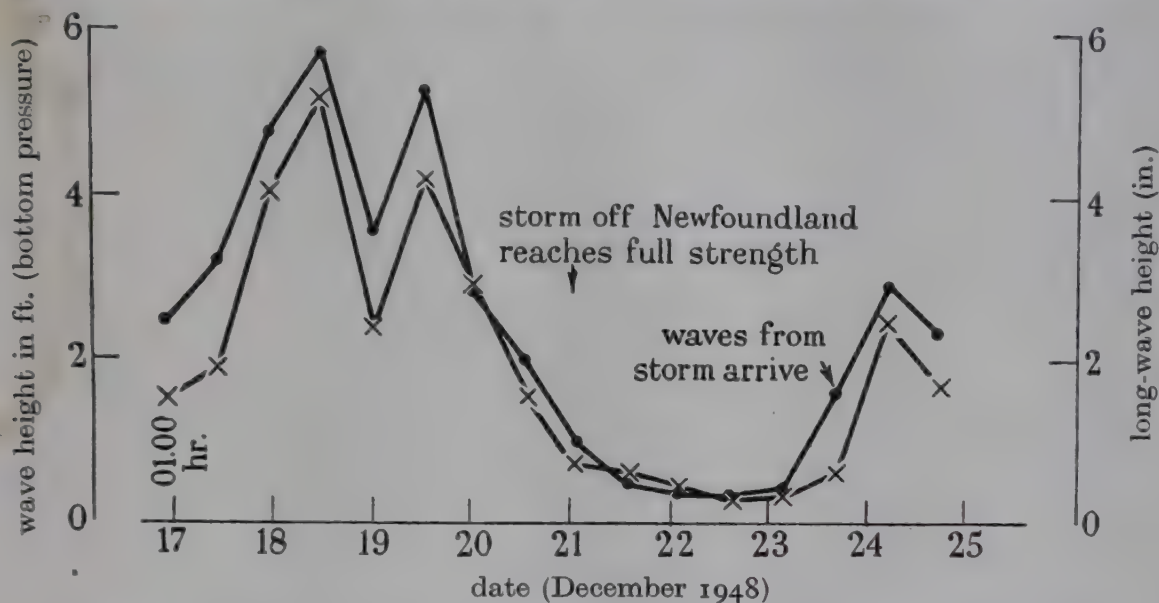


FIGURE 3. Variation of the ordinary-wave height and long-wave height with time: ●, ordinary waves; ×, long waves.

A detailed examination of the records (those in figure 7 for example) suggests that high groups of ordinary waves are followed by long waves after a time delay of 4 to 5 min. This is approximately the time taken for the groups of waves to reach the beach and for the long waves to travel back to the pressure recorder, which indicates that the long waves are produced by groups of high waves breaking on the beach. If this is so, the long-wave records should be similar to the envelope of the ordinary waves recorded 4 to 5 min. before.

To make a more thorough investigation, a correlation meter was constructed. It allows a comparison to be made between the long-wave record and the envelope of the ordinary waves, with the long-wave record delayed or advanced by varying amounts with respect to the ordinary-wave record. The meter produces a graph of the coefficient of correlation against the time displacement as one record is moved relative to the other. The apparatus and method are described in detail in appendix 1. Eight pairs of long-wave and short-wave-envelope records were compared, and the results are shown in figure 5. In the first five records at least, there is a marked correlation pattern when the comparison is made with the long-wave record lagging 4 to 5 min. behind the ordinary-wave envelope. The mean of the first five graphs,

shown in figure 6, strengthens this conclusion. A positive correlation means that an increase in wave height corresponds to an increase in water level. The shape of the correlation pattern gives an indication of the wave-form of the long wave produced by a single group of waves (or a single wave), and shows that such a group would propagate a disturbance which consists in sequence of a rise, fall and rise in water level. In pursuing the subject it must be remembered that the shape of the long waves may have been smoothed to some extent by the long-wave filter.

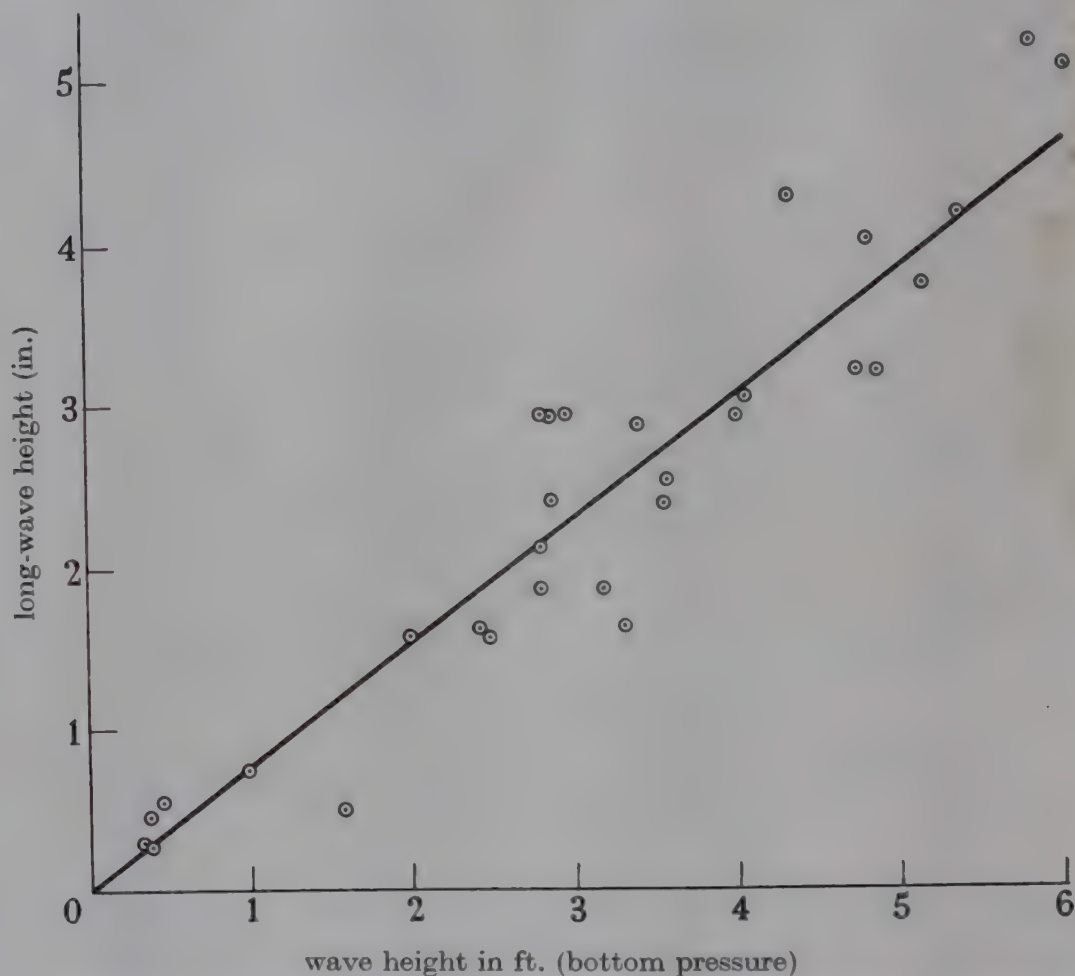


FIGURE 4. Relation between long-wave height and ordinary-wave height

In using the aneroid recorder with only a small volume of air in the rubber bag there is some danger that apparent 'long waves' might be due to an unsymmetrical response of the unit to positive and negative pressures when a group of high waves passes over it. The fact that there is only a small amount of correlation when the two records are compared with no time separation proves that this fear is unfounded.

A comparison of the ratio of long- to ordinary-wave heights with the period of the ordinary waves showed that the ratio was independent of the period of the waves.

CONCLUSION

If the evidence of figures 4 and 6, that the long waves are produced by the breaking of groups of high waves on the beach, is accepted, it becomes necessary to explain why the correlation between the envelope of the wave record and the long-



FIGURE 5. Change in correlation between the long waves and the ordinary-wave envelope with time separation of the records.

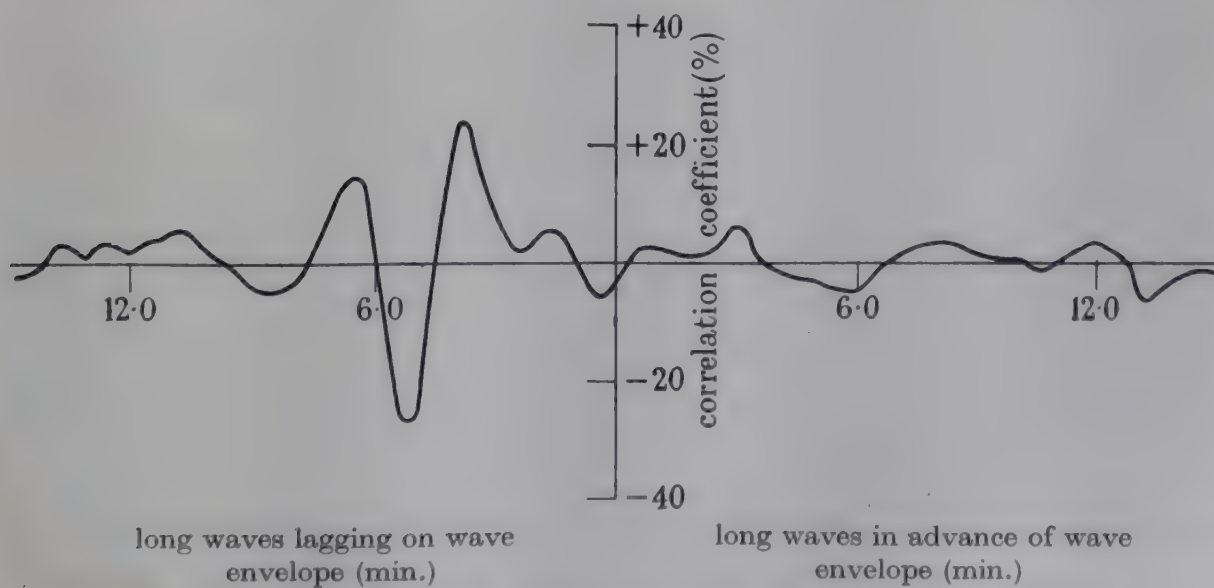


FIGURE 6. Mean of the five best correlograms (the first five in figure 5).

wave record delayed 4 to 5 min. is much better in some records than in others. The best correlation was obtained when the waves were produced by a distant storm. Under such conditions the wave groups were long and regular and would be expected to reach the beach without much change of shape; also, the length and distribution of the groups is such that the long waves from one group would not tend to be cancelled by those from adjacent groups. The worst correlation was

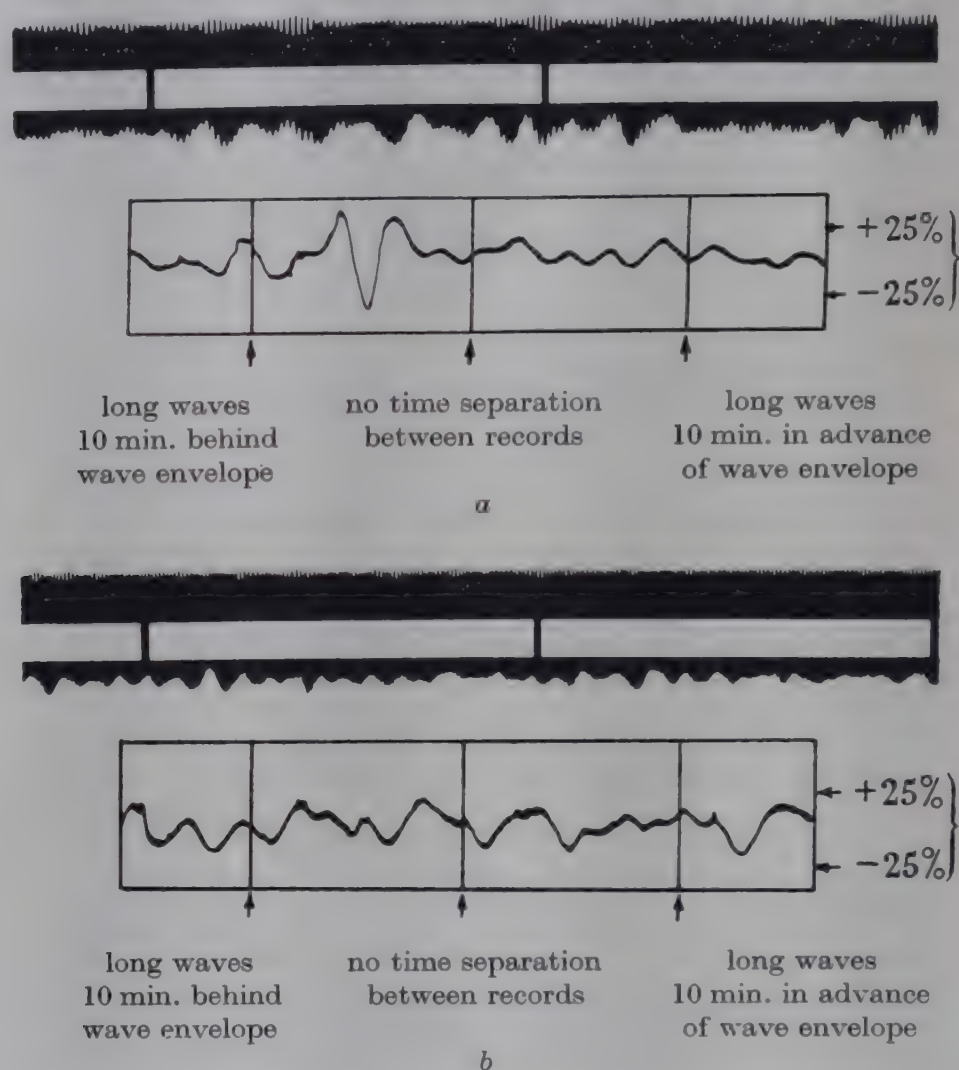


FIGURE 7. Wave record, long-wave record and correlation diagram for the best and worst correlation patterns: *a*, 05.00 hr. 24 December 1948; *b*, 22.00 hr. 16 December 1948.

obtained when the waves were generated over a large area near the coast and the groups were short and irregular. Such groups might change shape considerably before reaching the beach, and long waves from adjacent groups would tend to mask each other (see figure 7).

It seems reasonable at first to attribute the formation of the 'long waves' to the varying transport of the water in groups of high and low waves approaching the coast: the extra water released when a high group reaches the beach starts a surge, or a long wave, travelling out to sea. Such a simple explanation disagrees with the observations in two major respects: according to theory, the mass transport with a wave (in a given depth) is proportional to the square of the height, whereas the

observations show that the long-wave height is approximately linearly proportional to the ordinary-wave height. The simple explanation also requires that the long wave should be an elevation, whereas figure 6 shows that the outstanding feature of the observed wave is a depression in water level.

It is probably more correct to attribute the long wave to a sharp increase in shoreward transport of water in or near the breaking zone. In order to increase the forward mass transport, water is accelerated towards the beach, and there must be a corresponding acceleration out to sea because the momentum of the water moving in the two directions must be equal. To put it simply, the waves must have something to push against to accelerate water shorewards, and as a wave group reaches the breaker zone one long-wave elevation should be sent shorewards and another seawards, leaving a depression in between. The shorewards elevation would be reflected from the beach, and the observed sequence of elevation, depression and elevation would be established.

The increase in mass transport in the breaker zone may explain why the long-wave height is approximately linearly proportional to the height of the ordinary waves instead of to the square of their height, as might be expected from the theoretical result that the mass transport in a given depth of water is proportional to the square of the wave height. Low waves should undergo a proportionately greater increase in height and increase in mass transport than high waves because they travel farther up the beach before breaking; such an increase would make the transport in the low waves greater than the value expected from the square-law, and more in agreement with a linear relationship.

Considerable theoretical work has been done on the mass transport in waves, but there is still doubt as to what happens under actual conditions.

It is interesting to note that frequency analysis of a long-wave record shows that the energy of these 'beat-of-surf' waves lies almost entirely in the period range of 1 to 5 min.—a result that would be expected from the shape of the correlogram in figure 6. To detect still longer waves, such as may travel direct from a storm, it seems advisable to use an apparatus sensitive only to periods greater than 10 min.

APPENDIX 1. THE CORRELATION METER

A sketch of the apparatus is reproduced in figure 8. The two records to be correlated are prepared in the form of black and white profiles (see figure 7), and they are stretched round the outside of a celluloid sheet fixed in an arc of a circle centred on a gramophone turn-table. The gramophone turn-table carries two cylindrical lenses, and screens which cut off all light that does not pass through the lenses. The lamp filament is focused to a vertical line of light across the two records, and the rotation of the turn-table causes this to sweep from right to left along the records; two lenses are used to reduce the time during which no signal is picked up. The instantaneous value of the scattered light is proportional to the sum of the two ordinates of the records measured from their means plus an amount due to the mean width of white paper exposed; this unwanted mean value is

removed in the subsequent electronic circuit and may therefore be ignored. The scattered light is picked up on a photoelectric cell, the output of which is amplified and fed into a square-law meter using a thermistor bridge circuit which has a slow response and which gives an output proportional to the mean square of the signal. A mechanical device moves one record slowly along its time axis.

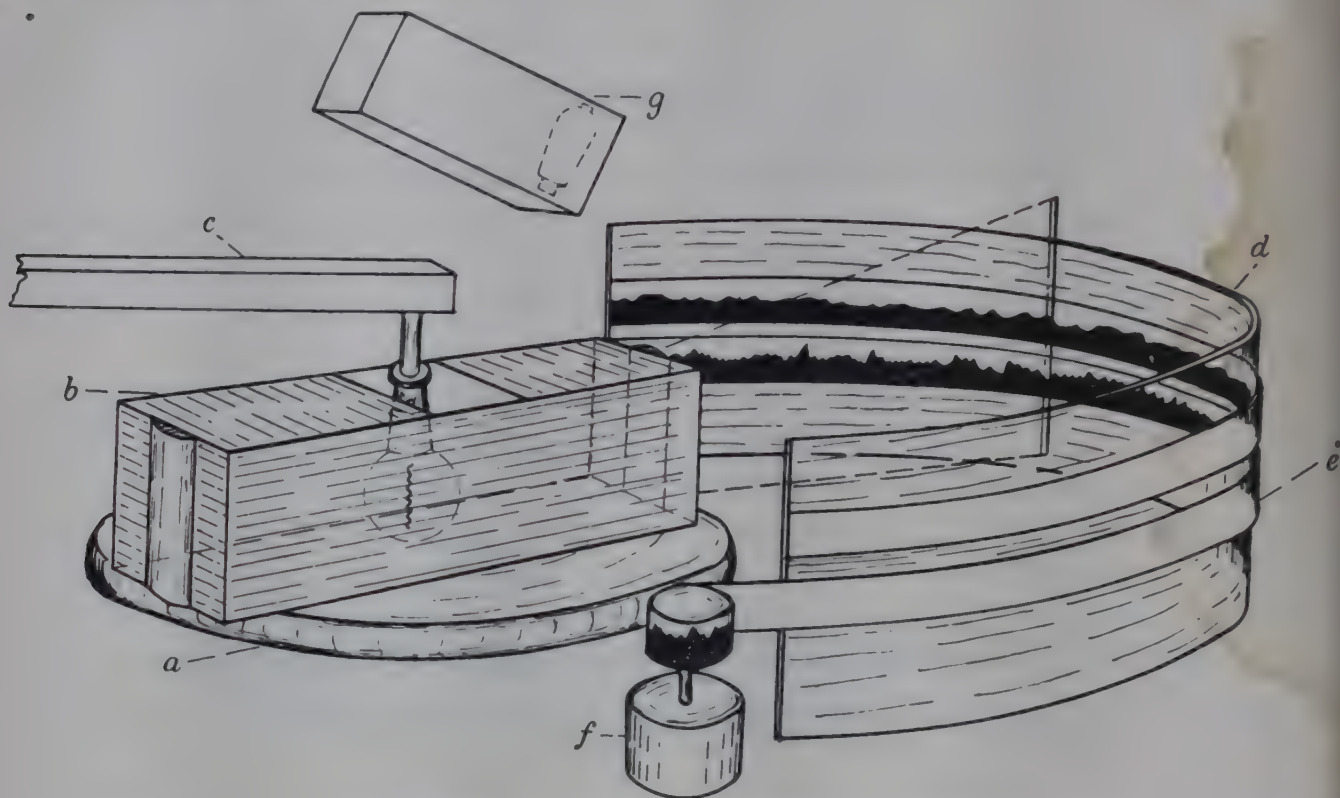


FIGURE 8. The correlation meter. *a*, gramophone turn-table carrying *b*, screened box with cylindrical lenses; *c*, stationary support for lamp; *d*, celluloid sheet; *e*, variable area records; *f*, device for slowly moving one record; *g*, light-sensitive cell.

It can be shown that the correlation coefficient r is given by

$$r = \frac{E_s - (E_A + E_B)}{2\sqrt{(E_A E_B)}},$$

where E_A is the mean square output from the meter with one record alone, E_B with the other record alone, and E_s with both records taken together.

If we assume that the mean square of the examined portion of the moving record remains constant, variations in the output of the square-law meter are proportional to variations in the correlation coefficient as the time separation of the record changes. The zero may be found, and the scale calibrated, by measuring the output when each of the records is on the meter alone.

To obtain records suitable for use on this machine without destroying the originals, the wave envelope on the original was blacked in with soft lead pencil, and the whole record printed by placing it face downwards on to photographic paper. It was covered with a sheet of Perspex held down by weights, and a bright light was then shone through for a suitable time. The print was processed, and a grid consisting of a series of lines at 1 in. intervals and perpendicular to the time

axis was drawn on the back. The long-wave record and the wave envelope were then cut off as long narrow strips, using a sharp knife and a steel straight edge, and in this form were ready for use.

In conclusion the author wishes to thank Mr N. F. Barber for his stimulating interest and helpful advice during the progress of this work, and the Admiralty for permission to publish the work.

REFERENCE

Munk, W. H. 1948 *Rev. Sci. Instrum.* **19**, 654.

Programme organization and initial orders for the EDSAC

BY D. J. WHEELER

University Mathematical Laboratory, University of Cambridge

(Communicated by D. R. Hartree, F.R.S.—Received 4 April 1950)

Orders for the machine are presented to it in coded form on punched tape, and are translated by the machine itself into the form in which they are held in the store, in accordance with a standard set of 'initial orders'. In the course of this input of orders it is possible to use the machine to make systematic modifications to the orders as punched on the tape; the modifications required, if any, are indicated by a code letter included in each order as punched on the tape. This system allows the use of a more flexible system of representing orders than the binary form used inside the machine.

It also enables sub-routines to be drawn up in a form independent of the particular values of parameters in them, the values of these parameters in any particular application of the sub-routine being inserted by the machine in the course of the initial input of orders, or in the course of operation of the machine. This simplifies the formation of a library of sub-routines and its use in the construction of long programmes.

A complete example is worked.

1. INTRODUCTION

This paper is concerned with some aspects of the organization of work for one kind of automatic calculating machine, of which the EDSAC, installed at the Mathematical Laboratory of the University of Cambridge, is an example. A general description of this machine has been given by Wilkes & Renwick (1949, 1950).

The essential parts of any automatic digital calculating machine are a store for holding numbers, an arithmetical unit, and a control unit. The store holds numbers in storage locations which are numbered so that we can refer to a storage location by means of an integer, known as the 'address' of that location, or of its content.

The essential features of the EDSAC for the purpose of this paper are three. First, it uses a single-address form for orders, that is, each order refers to a single storage location; the order code is given in appendix I. Secondly, the orders are

expressed inside the machine in coded form as numbers, and are contained in the same store as numbers, so that it is possible to modify orders by means of the arithmetic unit. Thirdly, the orders are normally placed in the store in such a way that the sequence of their addresses is the same as the time sequence in which they are to be carried out.

Another feature of the EDSAC, which affects some of the details but not the principles of this paper, is provision for grouping two adjacent numbers together and operating directly on the long number so formed. The content of the short-storage location n will be written $C(n)$ and that of the long-storage location m as $C(m')$; the content of the accumulator will be written as $C(Acc)$.

Some elementary aspects of programming for a machine with these features have already been discussed in a paper by Wilkes (1949).

2. INITIAL ORDERS

The orders which control the successive stages of the calculation, and the numerical data required in the course of it, are taken into the EDSAC from five-hole teleprinter tape. This tape is prepared using a keyboard perforator. This has been arranged so that each key corresponding to an integer between 0 and 9 punches a row of holes that represents the binary form of that number. The basic input operation is to read one row of holes on the tape and place the corresponding five-digit binary number in the store. This operation is controlled directly by orders held in the store, and, consequently, before the tape can be read it is necessary to place a group of orders in the store by some independent means. These orders are known as the initial orders and are permanently wired on a set of uniselectors. When the starting button is pressed, these orders are automatically transferred to the store and they then cause the input tape to be read. This means that the structure of the initial orders determines the way in which orders have to be punched on the tape. There are many forms the initial orders can take, and they should be chosen so that the work of preparing a tape is as simple as possible.

3. REPRESENTATION OF ORDERS

Inside the machine each order is expressed as a seventeen-digit binary number according to the following code. The five most significant digits represent the function of the order—add, subtract, etc. The sixth digit is '0'. The next ten digits give the address, m , of the storage location to which the order refers. The seventeenth digit is '0' if the order refers to a short number, and '1' if it refers to a long number. Since the first digit is regarded as a sign digit and a 'binary point' supposed to exist before the second digit, the number equivalent to the address m of an order is $m \times 2^{-15}$.

It is desirable that orders should be punched on the tape in the same form as they are written, without any modification or transliteration being necessary. The form in which orders are written and punched can be, and should be, chosen primarily with a view to the convenience of the user; the necessary conversion to the form in

which they are held in the store, including conversion of the address to the binary form, can then be arranged by a suitably chosen set of initial orders.

Two ways of writing orders have been used in connexion with the EDSAC, corresponding to two sets of initial orders. The second set of initial orders was introduced because the first set was found to suffer from certain disadvantages. Both sets will be described, since the first form of initial orders illustrate in a simple form some features also incorporated in the second set.

4. INITIAL ORDERS—FIRST FORM

With the first form of initial orders each order is written and punched as follows: first, a letter representing the five function digits—*A* corresponding to add, *S* to subtract, etc.; secondly, the address in decimal form with non-significant zeroes omitted (e.g. 15, not 0015); and thirdly an *S* or *L* according as the order refers to a short or a long storage location. For example, *A 97 S* stands for an order which causes the contents of 97 to be added to the accumulator, and *S 100 L* stands for an order which causes the contents of long storage location 100 to be subtracted from the accumulator. The abbreviated form of words '*A 97 S* adds *C* (97) to the accumulator' and '*S 100 L* subtracts *C* (100)' from the accumulator' will be used for these operations. The characters *S* and *L* are punched by the keyboard perforator in the binary forms of the decimal numbers 12 and 25, and can therefore be distinguished from decimal digits which are all less than 10.

The initial orders read the tape and assemble each order thereon in the following way:

The first row of holes—which corresponds to the function digits—is read, and the resulting binary digits are shifted to their correct digital position and placed temporarily in storage location 0. Subsequent rows of holes are then read. These will represent either decimal digits forming part of the address, or the characters *S* or *L* which terminates the address. The decimal digits are treated in sequence, a new partial address being formed each time, according to the following rules:

The first partial address is zero.

The $(n+1)$ th partial address equals ten times the n th partial address, plus the n th decimal digit of the address. Since this is done inside the machine the address is formed as a binary number.

For example, suppose the address is 594. The numbers will be written in their decimal form, but it must be remembered that the machine deals with them in their binary form, the transformation of single decimal digits into binary form being done by the keyboard perforator. After the function digits are read the first partial address is 0. This is multiplied by 10 and 5 added, giving the second partial address, 5. This is multiplied by 10 and 9 added, giving the third partial address, 59. This is multiplied by 10 and 4 added, giving the fourth partial address, 594.

The end of the address is identified by the character *S* or *L*. The order is then assembled from the function digits and the address, the length discriminant digit being made '0' or '1' according as the order is terminated by *S* or *L*.

The assembled order is then placed in the store; the first order goes into location 31, the second into location 32, and so on. When the last order has been taken in

from the tape, the programme is started by transferring control to the order stored in location 31.

Since the number of orders is different for each programme, it is necessary to punch on to the tape some indication which will cause control to be switched to the beginning of the programme when all the orders have been taken in. With the set of initial orders now being discussed, this is done by making $T(n+1)S$ the first order on the tape, where n is the storage location into which the last order is to be placed. The purpose of this order is to serve as a basis of comparison with an order which is formed in the course of the input process so that it can be determined when the input of the orders from the tape is complete.

The initial orders are given in detail in appendix II, and their action can there be followed in more detail with the aid of the notes given.

5. USE OF SUB-ROUTINES

Programmes are usually constructed from sequences of orders called sub-routines, each of which performs one stage in the solution of the problem. As many programmes have similar sub-routines, it is convenient to have a 'library' containing sub-routines for performing such standard operations as the evaluation of a sine, or a scalar product. The preparation of long programmes can then be simplified by making use of sub-routines from the library. These sub-routines, having been previously tested, will be free from errors, so that the chance of an accidental error occurring in a long programme will be minimized. Usually it will still be necessary to construct sub-routines for performing operations which are special to the particular problem in hand.

It is not convenient, however, to use sub-routines in combination with the initial orders described above. The difficulty is best illustrated by means of an example, and we will consider a sub-routine for finding the modulus of a given number. Such a short sequence of orders scarcely warrants being called a sub-routine, and would usually be incorporated directly into a programme. It will serve, however, as an example:

Example. Place the modulus of the contents of storage location 1 in storage location 10.

location of order in store	order	interpretation of order	general sequence of events
100	A 1 S	add $C(1)$ to the accumulator	add contents of 1 to the accumu- lator
101	E 104 S	test sign of the number in the accumulator; if positive, proceed with order 104; if negative, proceed with order 102	change sign of the number in accu- mulator if it is negative
102	T 10 S	transfer $C(Acc)$ to 10 and clear accumulator	
103	S 10 S	subtract $C(10)$ from the accumulator	
104	T 10 S	transfer $C(Acc)$ to location 10 and clear the accumulator	transfer to location 10

It can easily be seen that if these orders had been placed in the store from location m onwards instead of from 100, the second order would have been $E(m+4)S$, i.e. the address specified in this order depends on the location in which the sub-routine is to be placed, whereas the addresses in the other orders do not. If the initial orders given above were used each time the sub-routine was used, it would be necessary before punching the tape to adjust appropriately all orders referring to addresses which are dependent on the location of the sub-routine.

The adjusting process follows fixed rules, and therefore the machine itself can be used to do it. The purpose of the second form of initial orders was to accomplish this in the most convenient way. As a result it is now possible to store sub-routines as strips of tape and to copy them mechanically on to any programme tape which requires them.

6. PARAMETERS

By a parameter of a sub-routine is meant a number (often an address referred to in the course of the sub-routine) which may have different values in different applications of the sub-routine; for example, if we had a sub-routine for placing $|C(r)|$ in storage location s , for any values of r and s and not only for the values $r = 1, s = 10$ of the example already considered, r and s would be parameters of the sub-routine. It is clearly economical and convenient to be able to draft, and to punch, sub-routines in such a form that a single sub-routine can be used for different values of the parameters. For example, it is more economical to have a single sub-routine capable of finding the scalar product of n elements where n can be made to take any value from 1 to 20 as desired, rather than twenty different sub-routines.

It is desirable to distinguish between two types of parameters, pre-set and programme parameters. A programme parameter is one which is used where it is necessary to determine or to adjust its value *during the execution of a programme*; its value is inserted into the sub-routine by the machine in the course of the calculation. A pre-set parameter, which arises only from the use of a library, is one whose value is known and constant throughout the execution of any one programme, although it may vary from one programme to another. The second form of initial orders have been chosen so that the orders are adjusted according to the values of these pre-set parameters during the input of orders. The use of pre-set parameters has the advantage that the programme need not include provision for making these adjustments.

7. INITIAL ORDERS—SECOND FORM

These orders are now wired on the uniselectors in place of the previous ones. Orders are punched and written in the following way: first a letter to represent the function; digits; secondly, the address in decimal form; and thirdly, a code letter. The code letter does not strictly form part of the order but controls the action of the initial orders.

In addition to orders it is possible to punch 'control combinations' on the tape. These have a form similar to that of orders, but are used to control and direct the action of the initial orders. They are as follows:

- T m K* causes the orders which follow to be placed in locations $m, m+1, m+2$, and so on. This can be used to place groups of orders where they are needed in the store.
- G K* is punched before a sub-routine and causes the address of the storage location into which the first order of the sub-routine is put to be placed in storage location 42.
- T Z* causes the order which follows to be placed in storage location p , and subsequent orders in $p+1, p+2$, etc., where p is the address held in 42.
- E n KPF* causes control to be transferred to location n and so starts the programme.

The method of starting the programme used here is slightly more convenient in practice than the method used with the previous set of initial orders.

The function of the code letters is to specify the location of numbers which are to be added to the orders as punched on the tape before they are placed in the store. The code letters are as follows:

code letter	integer equivalent to code letter	storage location containing number specified	value of number
<i>F</i>	17	41	zero
θ	18	42	$m \cdot 2^{-15}$. Address of first order of sub-routine
<i>D</i>	19	43	2^{-16}
ϕ	20	44	these can be set to any value from the tape using control combinations, e.g. $T\ 45\ K, P\ 10\ F^\dagger, P\ 20\ F$ cause $P\ 10\ F$ to be put in 45 and $P\ 20\ F$ to be put in 46, etc.
<i>H</i>	21	45	
<i>N</i>	22	46	
<i>M</i>	23	47	
Δ	24	48	
<i>L</i>	25	49	
<i>X</i>	26	50	
<i>G</i>	27	51	
<i>A</i>	28	52	
<i>B</i>	29	53	
<i>C</i>	30	54	
<i>V</i>	31	55	

F and *D* can be considered as corresponding directly to the *S* and *L* of the previous initial orders. If an order of a sub-routine has the code letter θ , the address specified in it will be increased by the number held in location 42 which is the address of the first order of that sub-routine. If an order has code letter *H*, *N*, *M*, etc., the address specified in it will be increased by the value of the content of the storage location selected by that code letter, as given by the table.

There is one additional code letter, π . This is only used in combination with one of the other code letters and is always placed directly before that code letter. It causes the associated order to refer to a long number.

These initial orders operate in the following way:

(1) The function digits are read from the tape and shifted into their correct digital position and stored temporarily in storage location 40.

$\dagger$ The character *P* is equivalent to five zeroes, so that $P\ 10\ F$ is numerically equal to 10×2^{-15} .

(2) The address is formed by the same method as that used with the first form of initial orders, the end of the address being indicated by a code letter which has a numerical value greater than ten.

(3a) The code letter is read and tested to see whether its numerical value is greater than 16; if so, it is then used to form an order which selects the appropriate number to be added to the address of the order that is being formed. The modified address and function digits are assembled and placed in the store.

(3b) If the code letter has a numerical value less than 17, control is switched to the appropriate sequence of operations defined by the code letter, e.g. K , Z , π .

This process continues until a control combination is encountered when the process is adjusted accordingly. The action of these initial input orders can be followed in greater detail in appendix III.

It will be seen that with this method of forming an order from the characters punched on the tape, non-significant zeroes in the address specified in an order do not have to be punched; and further that for this purpose the address of location 0 can be treated as a non-significant zero. It is therefore omitted in writing orders; for example, $T F$ is used for $T 0 F$.

8. USE OF SUB-ROUTINES WITH THE SECOND FORM OF INITIAL ORDERS

The sub-routine previously referred to for taking the modulus of a number is given below in the notation corresponding to the second form of initial orders. The action of the machine during the input of the orders is as follows; it is supposed that they are being placed in storage locations n , $n+1$, $n+2$, etc.

entry on the input tape	action during input
$G K$	this control combination puts the number $n \times 2^{-15}$ in storage location 42
$A 1 F$	this is placed in location n
$E 4 \theta$	$n \times 2^{-15}$ is added and the result $E(n+4) F$ is placed in location $n+1$
$T 10 F$	these are not modified but are placed in locations $n+2$, $n+3$ and $n+4$
$S 10 F$	
$T 10 F$	

Before being placed in the store these orders are converted to their binary form. We represent the orders using the new notation, code letters F and D corresponding to the values 0 and 1 respectively of the length-discriminant digit:

location	orders
n	$A 1 F$
$n+1$	$E n+4 F$
$n+2$	$T 10 F$
$n+3$	$S 10 F$
$n+4$	$T 10 F$

It will be seen that these orders have now been adjusted with reference to the location in which they were to be placed, and have been placed there. By the use of a similar example we can show how we make use of the code letters H , N , ... V . Suppose we need to place the modulus of the number in long storage location 200 in long storage location 100. It is unlikely that we would have such a sub-routine, but we may have a general sub-routine in the library, 'place the modulus

of the contents of the location defined by the parameter in 45 (code letter *H*) in the location defined by the parameter in 46 (code letter *N*)'. Such a library tape is given below:

T	Z	control combination	} library tape
A	H		
E	4θ	orders	
T	N		
S	N		
T	N		

GK, T 45 K, P 200 D, P 100 D, is punched on the input tape before the library tape is mechanically copied on to it. Then the input tape becomes

<i>G</i>	<i>K</i>	} control combination
<i>T</i>	<i>45 K</i>	
<i>P</i>	<i>200 D</i>	} punched pre-set parameters
<i>P</i>	<i>100 D</i>	
<i>T</i>	<i>Z</i>	} copied from library tape
<i>A</i>	<i>H</i>	
<i>E</i>	<i>4 θ</i>	
<i>T</i>	<i>N</i>	
<i>S</i>	<i>N</i>	
<i>T</i>	<i>N</i>	

The action during the input of orders is as follows:

<i>G</i>	<i>K</i>	puts <i>P n F</i> in 42, where <i>n</i> is the location of the first order of the sub-routine
<i>T</i>	<i>45 K</i>	arranges for parameters to be put in 45, 46, ...
<i>P</i>	<i>200 D</i>	is put in 45) in the machine <i>P</i> corresponds to zero so that <i>P 200 D</i> can be con-
<i>P</i>	<i>100 D</i>	is put in 46) sidered as just an address to be inserted in an order
<i>T</i>	<i>Z</i>	arranges for the following orders to be put in <i>n, n + 1</i> , etc. (<i>n</i> is in 42).
<i>A</i>	<i>H</i>	order becomes (<i>A F</i>) + (<i>P 200 D</i>), i.e. <i>A 200 D</i> , and is placed in <i>n</i> .
<i>E</i>	<i>4 θ</i>	order becomes (<i>E 4 F</i>) + (<i>P n F</i>), i.e. <i>E n + 4 F</i> , and is placed in <i>n + 1</i> .
<i>T</i>	<i>N</i>	order becomes (<i>T F</i>) + (<i>P 100 D</i>), i.e. <i>T 100 D</i> , and is placed in <i>n + 2</i> and so on.

The orders are thus put in the store in the desired form.

9. THE USE OF SUB-ROUTINES IN CONSTRUCTING PROGRAMMES

The simplest way of using sub-routines to form a programme is to arrange them in sequence in the store so that when one sub-routine is finished, the next one is commenced, and so on. Sub-routines used in this way will be called *open* sub-routines.

This method of arranging sub-routines becomes inconvenient if it is required to use the same sub-routine at more than one point of the programme; it is then necessary to switch control to and from the sub-routine at the appropriate time. In these circumstances it is more convenient to use a sub-routine into which the necessary switching order for transfer of control on completion of the sub-routine can be incorporated, so that it can be used from any point of the store, control being returned—when the calculation is completed—to a point immediately after the point from which the sub-routine was called in. This will be called a *closed*

sub-routine. There are various methods of arranging the operation of entering a closed sub-routine and returning to the main programme after the operation of the sub-routine has been completed. The method used with the EDSAC is given below; m is the address of the first order of the sub-routine.

location	order			notes
				orders calling in the sub-routine
n	A	n	F	accumulator assumed empty
$n+1$	G	m	F	this order adds itself into the accumulator control to m since $C(Acc)$ is negative, as the numerical equivalent of the character A is negative
$n+2$				programme after sub-routine
				sub-routine
m	A	3	F	accumulator contains $A n F$ adds $U 2 F$ to the accumulator forming $(U + A)(2 + n) F$, numerically equal to the order $E n + 2 F$
$m+1$	T	$m+r$	F	transfers $E (n + 2) F$ to the operational end of the sub- routine where it is used to return control to the main programme
...				accumulator empty
$m+r$	$(E$		$F)$	becomes $E n + 2 F$ as a result of the operation of the order in $(m+1)$. When $(m+r)$ is reached in the sequence of operations of the sub-routine, control is switched to $(n+2)$

It should be noted that all orders in this example have been written with their true addresses, that is, in the form they take in the store, whereas on the tape they would have been punched in terms of θ , etc.

10. USE OF PARAMETERS IN CLOSED SUB-ROUTINES

As mentioned before, parameters associated with sub-routines greatly extend their scope. There are different places in which programme parameters can be stored, but one of the more convenient places in which to store them is in the main programme immediately following the orders calling in the sub-routine. Arranging for the parameters to be stored in the main programme leads to economy of thought, as the necessary values of the parameter to be used are considered with the 'calling-in' orders. Control is, of course, returned to the order in the main programme which follows the parameters. For example, a closed print sub-routine which prints $C(0')$ to m figures, where $P m F$ is a programme parameter specified in the main programme, would be 'called in' as follows:

location			
p	A	p	F
$p+1$	G	s	F
$p+2$	P	m	F
			programme parameter

s is the address of the start of the print routine.

This sub-routine can find and use the parameter $P m F$ because the address specified in the order held in the accumulator on entering the sub-routine provides a reference number.

Suppose the sub-routine has been drawn up in such a way that the parameter $P m F$ is used by being added into the accumulator during the execution of the

sub-routine. Since the parameter is in location $(p+2)$, the order required for this is $A(p+2)F$. The 'link' order required to return control to the main programme is $E(p+3)F$. These orders can be formed as follows:

To $A(p)F$ contained in the accumulator on entering the sub-routine we add $A(p+2)F - A(p)F$ and so obtain $A(p+2)F$. This is placed in the appropriate position in the sub-routine.

To $A(p+2)F$ we add $E(p+3)F - A(p+2)F$ in order to obtain the link order $E(p+3)F$.

$$\begin{aligned} \text{Now } A(p+2)F - A(p)F &= (A - A)2F \\ &= P2F \text{ as } P \text{ corresponds to zero} \end{aligned}$$

$$\begin{aligned} \text{and } E(p+3)F - A(p+2)F &= (E - A)1F \\ &= U1F \text{ (cf. §9).} \end{aligned}$$

These two constants are independent of p and would be included in the sub-routine.

11. EXAMPLE OF CODING USING LIBRARY SUB-ROUTINES

An example will now be given of the way in which a complete calculation could be coded. Suppose it is required to tabulate ψ , where

$$\begin{aligned} \cos \psi &= \sin \omega \cos \delta + \cos \omega \sin \delta \sin \phi \\ &= \sin(\omega + \delta) - (1 - \sin \phi) \cos \omega \sin \delta; \end{aligned}$$

δ varies from 0° to 23° in 1° steps; ϕ varies from 0° to 90° in 1° steps; ω is a given constant. ψ is to be tabulated in degrees and decimals.

There are various ways in which this calculation could be done; the one which is used here as an illustration is as follows:

$\sin \delta$ is calculated;
 $\sin \delta \cos \omega$ is calculated;
 $\sin \phi$ is calculated;
 $\sin \delta \cos \omega \sin \phi - \sin \delta \cos \omega$ is calculated;
 $\sin(\delta + \omega)$ is calculated;
 $\cos \psi$ is calculated;
 ψ is calculated in radians;
 ψ is multiplied by a factor to convert it into degrees and decimals, and printed;
the current value of δ is tested to find out if it has reached 23; if it has not, the current value of δ is increased by 1° and the calculation repeated for this new value of δ ; if it has, the current value of ϕ is tested. If this has reached 90° then the calculation is finished; if it has not then δ is made zero and the current value of ϕ is increased by 1° and the calculation repeated.

A list of the sub-routines required is drawn up. They are

- (a) decimal input sub-routine; this is required to convert the decimals to binary form and place them where they are required in the store;
- (b) sine sub-routine;
- (c) inverse cosine sub-routine;
- (d) print sub-routine; this converts the results to decimal form and prints them in a layout specified by values of pre-set parameters.

The index of the library is consulted and the following sub-routines selected:

- (a) Sub-routine $R 5$. This is a sub-routine of special form, not conforming to either the open or the closed form. It converts decimal numbers on the input

tape into the binary form required for the machine. It is used before the main programme has been taken into the store so that the space occupied by it may be used again for other sub-routines. It is followed on the tape first by a parameter $T m D$, which specifies where the numbers are to be placed in the store, next by the numbers themselves punched in decimal form, and lastly by a control combination $X T Z$. This causes control to be returned to the initial orders so that more orders may be taken into the store.

(b) Sub-routine $T 2$. This is a closed sub-routine. It causes $\frac{1}{2} \sin [2 C (4')]$ to be placed in long storage location 4, $C (4')$ being an angular measure in radians. It occupies 42 storage locations, and its first order must be placed in an even location.

(c) Sub-routine $T 5$. This is a closed sub-routine. It causes $\frac{1}{2} \cos^{-1} [2 C (4')]$ (radians) to be placed in long storage location 4. It occupies 58 storage locations and its first order must be placed in an even storage location.

(d) Sub-routine $P 4$. This is a closed sub-routine. It causes the positive number (which must be less than 1) held in long storage location 0 to be printed to r decimals, the numbers being arranged in c columns with s spaces between them; r , c , and s are pre-set parameters of the sub-routine and the values to be given to them in any particular application of this sub-routine are specified as follows. Suppose 12 columns, two spaces between columns, and 4-figure entries are required. Then

$G \quad K$
 $T \ 45 \ K$
 $P \ 12 \ F$
 $P \ 2 \ F$
 $P \ 4 \ F$

is punched on the tape directly in front of this sub-routine. It occupies 55 storage locations.

12. THE CODING OF THE MAIN PROGRAMME

The next step is to plan and code the main programme. In order to do this, an indication is required for the locations to be used for numbers and constants. It is not convenient to fix the actual locations until later, and for this reason the code letter H will be used in orders referring to these numbers and constants. Then the actual locations may be fixed later by assigning the appropriate numerical value to the parameter m to be used in instructions with code letter H . The initial values of all these numbers are placed in their respective storage locations by the decimal input sub-routine, as follows:

location	number
$m + 0$	current value of $\frac{1}{2}\delta$
2	current value of $\frac{1}{2}\phi$
4	end value of $\frac{1}{2}\delta$, $11\frac{1}{2}^{\circ}$ *
6	end value of $\frac{1}{2}\phi$, 45° *
8	$\cos \omega$
10	$\frac{1}{2}\omega$
12	$\frac{180}{\pi} \times \frac{1}{100}$ a constant which converts radians to hundreds of degrees

* In practice the values stored here would be $11\frac{3}{4}^{\circ}$ and $45\frac{1}{4}^{\circ}$ to ensure the correct number of repetitions of the calculation independently of rounding errors.

Suppose that the starting location of the sine sub-routine is indicated by the code letter N , that of the inverse cosine sub-routine by M , and that of the print sub-routine by Δ . These values can also be fixed later. The location of the main programme in the store is also undetermined for the present; it is therefore coded just like a sub-routine, with the code letter θ for those orders in which the address has to be adjusted in accordance with the address into which the first order of the programme is placed. In the explanation, 'to 10' will mean 'transfers to long storage location 10'.

number of order in main programme	order	action of orders	general sequence of events
	T Z	control combination on the tape	
0	T F	clears the accumulator	
1	A H	$\frac{1}{2}\delta$ to 4'	calculates $\frac{1}{2} \sin \delta$
2	T 4 D		
3	A 3 θ		
4	G N		
5	H 4 D	calls in the sine routine; results in $\frac{1}{2} \sin \delta$ to 4'	
6	V 8 H		
7	Y F		
8	T 10 D		
9	A 2 H	$(\frac{1}{2} \sin \delta) \cos \theta$ to 10'	calculates $\frac{1}{2} \cos \theta \sin \delta$
10	T 4 D		
11	A 11 θ		
12	G N		
13	H 4 D	$\frac{1}{2}\phi$ to 4'	calculates $\frac{1}{2} \sin \phi$
14	V 10 D		
15	L D		
16	S 10 D		
17	Y F	$\frac{1}{2} \sin \phi$ to 4'	
18	T 10 D		
19	A 10 H		
20	A H		
21	T 4 D	$\frac{1}{2} (\omega + \delta)$ to 4'	calculates $\frac{1}{2} \cos \psi$
22	A 22 θ		
23	G N		
24	A 4 D		
25	A 10 D	$\frac{1}{2} \cos \psi$ to 4'	calculates $\frac{1}{2} \psi$
26	T 4 D		
27	A 27 θ		
28	G M		
29	H 4 D	calls in inverse cosine routine and results in $\frac{1}{2} \psi$ to 4'	
30	V 12 H		
31	L D		
32	Y F		
33	T D	expresses the angle in units of 100° , so that its measure is less than 1. The first two digits of the number printed will represent the integral number of degrees in the angle, the other digits re- present the fraction of a degree contained in that angle	
34	A 34 θ		
35	G Δ		
36	A H		
37	S 4 H	tests to see if the current value of δ has reached $11\frac{1}{2}^\circ$. If it has, switches control to 44 θ	
38	E 44 θ		
39	T F		
40	A H		
41	A 14 H	clears the accumulator	
42	T 14 H		
		increases $\frac{1}{2}\delta$ by $\frac{1}{2}^\circ$	

number of order in main programme	order	action of orders
43	<i>E</i> θ	starts calculation again
44	<i>T</i> <i>F</i>	reached from 38 if current value of δ has reached $11\frac{1}{2}^\circ$. Clears accumulator
45	<i>A</i> 2 <i>H</i>	} tests if current value of ϕ has reached 90° ; if it has, switch control to 54
46	<i>S</i> 6 <i>H</i>	
47	<i>E</i> 54 θ	
48	<i>T</i> <i>F</i>	clears the accumulator
49	<i>A</i> 2 <i>H</i>	} increases $\frac{1}{2}\phi$ by $\frac{1}{2}^\circ$
50	<i>A</i> 14 <i>H</i>	
51	<i>T</i> 2 <i>H</i>	
52	<i>T</i> <i>H</i>	makes (current value of δ) = 0
53	<i>E</i> θ	re-starts calculation for next value of ϕ
54	<i>Z</i> <i>F</i>	stop order; reached from 47 if current value of ϕ is 90°

The next step is to decide on which positions the various sub-routines and numbers will occupy in the store. When this has been done the programme tape can be constructed.

Let the locations of the sub-routines be chosen as follows: Input sub-routine *R* 5 starts at location 60. Sine sub-routine *T* 2 starts at location 60 (this goes into the storage locations previously occupied by *R* 5). Since this occupies 42 storage locations it ends at 101. The inverse cosine sub-routine must start at an even location; let it start at 102; it then ends at 159. Print sub-routine *P* 4 starts at storage location 160 and ends at 214. The main programme starts at storage location 215 and ends at 269, and the numbers start at long storage location 270. Thus the parameter for the input routine must be *T* 270 *D* and the main programme must have the following symbols punched in front of it on the input tape:

G *K*
T 45 *K*
P 270 *D* used by code letter *H*
P 60 *F* used by code letter *N*
P 102 *F* used by code letter *M*
P 160 *F* used by code letter Δ

At the end of the main programme the control combination *E* 215 *K*, *P* *F* is required so that after taking in the whole programme the machine will return to the start of the main programme at address 215 to start the calculation itself.

The programme tape is then punched as follows:

storage locations to be used	tape entries	notes
	<i>T</i> 60 <i>K</i>	control combination to ensure that the first order of the first sub-routine is placed in 60
60*	input sub-routine <i>R</i> 5	copied from library tape
	<i>T</i> 270 <i>D</i>	parameter used by <i>R</i> 5
	numbers and constants	placed by <i>R</i> 5 in 270', 272', etc.
60 } 101 }	sine sub-routine <i>T</i> 2 }	copied from library tape
102 } 156 }	inverse cosine sub-routine <i>T</i> 5 }	copied from library tape

* Temporarily, while input of numbers is taking place.

storage
locations
to be used

tape entries

notes

		<i>G</i>	<i>K</i>	}	sets parameters for main programme
		<i>T</i>	45 <i>K</i>		
		<i>P</i>	270 <i>D</i>		
		<i>P</i>	60 <i>F</i>		
		<i>P</i>	102 <i>F</i>		
		<i>P</i>	160 <i>F</i>		
215	}	main programme			
269					
		<i>E</i>	215 <i>K</i>	}	control combination to start the programme
		<i>P</i>	<i>F</i>		

APPENDIX I. EDSAC ORDER CODE

<i>A</i>	<i>n</i>	Add the number in storage location <i>n</i> into the accumulator.
<i>S</i>	<i>n</i>	Subtract the number in storage location <i>n</i> from the accumulator.
<i>H</i>	<i>n</i>	Transfer the number in storage location <i>n</i> into the multiplier register.
<i>V</i>	<i>n</i>	Multiply the number in storage location <i>n</i> by the number in the multiplier register and add into the accumulator.
<i>N</i>	<i>n</i>	Multiply the number in storage location <i>n</i> by the number in the multiplier register and subtract from the accumulator.
<i>T</i>	<i>n</i>	Transfer the contents of the accumulator to storage location <i>n</i> and clear the accumulator.
<i>U</i>	<i>n</i>	Transfer the contents of the accumulator to storage location <i>n</i> and do not clear the accumulator.
<i>R</i>	$2^{n-2}F^*$	Shift the number in the accumulator <i>n</i> places to the right; i.e. multiply it by 2^{-n} .
<i>L</i>	$2^{n-2}F^*$	Shift the number in the accumulator <i>n</i> places to the left; i.e. multiply it by 2^n .
<i>E</i>	<i>n</i>	If the number in the accumulator is greater than or equal to zero execute next the order which stands in storage location <i>n</i> ; otherwise proceed serially.
<i>G</i>	<i>n</i>	If the number in the accumulator is less than zero execute next the order which stands in storage location <i>n</i> ; otherwise proceed serially.
<i>I</i>	<i>n</i>	Read the next row of holes on the tape and place the resulting 5 digits in the least significant places of storage location <i>n</i> .
<i>O</i>	<i>n</i>	Print the character now set up on the teleprinter and set up on the teleprinter the character represented by the five most significant digits in storage location <i>n</i> .
<i>F</i>	<i>n</i>	Place the five digits which represent the character next to be printed by the teleprinter in the five most significant places in storage location <i>n</i> .
<i>Y</i>		Round off the number in the accumulator to 34 binary digits.
<i>Z</i>		Stop the machine and ring the warning bell.

* For a shift of one place right or left the forms of the order are *RD*, *LD* respectively (with first form of initial orders, *F* and *D* are replaced by *S* and *L*).

APPENDIX II. INITIAL ORDERS 1

When the starting button is pressed these orders are placed in the store in locations 0 to 30 and control is set so that the first order obeyed is in location 0.

location of order	order	action of order	general sequence of events
0	<i>T S</i>	clears accumulator after transferring <i>C (Acc)</i> to location 0	
1	<i>H 2 S</i>	places <i>C (2)</i> in multiplier register. This <i>T S</i> is 10/32 numerically	clears accumulator and puts 10/32 in multiplier register
2	<i>T S</i>	clears accumulator	
3	<i>E 6 S</i>	examines <i>C (Acc)</i> . As this is zero, control is transferred to the order in location 6.	control switched to 6. These locations 0-3 are then used as 'working space'
4	<i>P 1 S</i>	2^{-15}	constants
5	<i>P 5 S</i>	10×2^{-15}	
6	<i>T S</i>	clears accumulator	
7	<i>I S</i>	input of function digits to location 0	input of function digits to their correct digital location in 0
8	<i>A S</i>	adds function digits into accumulator	
9	<i>R 16 S</i>	shifts <i>C (Acc)</i> six places to the right so that the function digits are in the five most significant places of the least significant half of the accumulator	
10	<i>T L</i>	transfers <i>C (Acc)</i> to long location 0 so the function digits are in their correct positions in short storage location 0	
11	<i>I 2 S</i>	input of next character on the tape to storage location 2	reads character on the tape and tests whether it is numerically less than 10
12	<i>A 2 S</i>	adds <i>C (2)</i> to <i>C (Acc)</i>	
13	<i>S 5 S</i>	subtracts <i>C (5)</i> from <i>C (Acc)</i>	
14	<i>E 21 S</i>	if <i>C (Acc) ≥ 0</i> , control is transferred to order in location 21	
15	<i>T 3 S</i>	clears the accumulator using 3 as a rubbish dump	one stage of the binary decimal conversion. New partial address is obtained by taking 10 times old partial address and adding the next digit
16	<i>V 1 S</i>	<i>C (1)</i> , multiplied by 10/32 which is held in multiplier register, is added to the accumulator	
17	<i>L 8 S</i>	shifts <i>C (Acc)</i> five places to the left, i.e. <i>C (Acc)</i> is multiplied by 32	
18	<i>A 2 S</i>	adds <i>C (2)</i> to <i>C (Acc)</i>	
19	<i>T 1 S</i>	transfers <i>C (Acc)</i> to storage location 1	control is transferred to 21 from the order 14 when character read is <i>S</i> or <i>L</i> . When <i>L</i> has been read the 17th digit of the accumulator is a 1, when <i>S</i> has been read it is a 0
20	<i>E 11 S</i>	As <i>C (Acc) = 0</i> , control is transferred to the order in location 11	
21	<i>R 4 S</i>	shifts <i>C (Acc)</i> four places to the right	
22	<i>A 1 S</i>	adds <i>C (1)</i> to <i>C (Acc)</i>	
23	<i>L L</i>	shifts <i>C (Acc)</i> one place to the left	the address has been formed $\times 2^{-15}$ and so needs doubling forms the complete order in accumulator
24	<i>A S</i>	adds <i>C (0)</i> to <i>C (Acc)</i>	
25	<i>T 31 S</i>	transfers <i>C (Acc)</i> to 31. The address specified in this order is increased by 1 each time the order is used and it places the constructed orders in sequence in the store	transfers the order to its final position in the store
26	<i>A 25 S</i>	adds <i>C (25)</i> to <i>C (Acc)</i>	increases the address specified in order 25 by 1; e.g. <i>T 31 S</i> is replaced by <i>T 32 S</i> and so on
27	<i>A 4 S</i>	adds <i>C (4)</i> to <i>C (Acc)</i>	
28	<i>U 25 S</i>	transfers <i>C (Acc)</i> to 25 without clearing accumulator	tests whether all orders have been taken in. Location 31 contains order <i>T (n+1) S</i> , the first order to be placed in the store; and so <i>C (Acc)</i> will be positive only when all orders have been taken into the store
29	<i>S 31 S</i>	subtracts <i>C (31)</i> from <i>C (Acc)</i>	
30	<i>G 6 S</i>	if <i>C (Acc) ≤ 0</i> , control is transferred to the order in location 6	
end of initial orders			
31	<i>T (n+1) S</i>		the first order to be placed in the store

APPENDIX III. INITIAL ORDERS—SECOND FORM

When the starting button is pressed these orders are placed in the store locations 0 to 40 and control is set so that the first order obeyed is in location 0. In the notes b stands for the numerical equivalent of the character last read from the tape.

location of order	order	action of order	general sequence of events
0	$T \quad F$	clears accumulator using 0 as a 'rubbish bin'	the accumulator is cleared and control is switched to 20; thereafter storage locations 0, 1 are used as working space constants intended to be left here unaltered for use in any programme
1	$E \ 20 \ F$	$C(Acc) = 0$ so that control is switched to order 20	
2	$P \ 1 \ F$	2^{-16}	
3	$U \ 2 \ F$	not used by initial orders	
4	$A \ 39 \ F$	adds 11×2^{-16} to $C(Acc)$	construction of the address. These orders are entered at 8 with $C(Acc) = 0$. The next character, represented by b on the input tape is read and tested to see if it is less than 11; if so it is doubled and added to ten times $C(0)$, the sum being sent back to 0. This cycle will be repeated until a code letter is read from the tape, when the sequence of orders starting at 13 is carried out
5	$R \ 4 \ F$	shifts $C(Acc)$ four places to the right, i.e. divides $C(Acc)$ by 16	
6	$V \quad F$	multiplies $C(0)$ by contents of multiplier register which is holding 10/32	
7	$L \ 8 \ F$	shifts $C(Acc)$ five places to the left, i.e. multiplies $C(Acc)$ by 32	
8	$T \quad F$	transfers $C(Acc)$ to 0 and clears Acc	
9	$I \ 1 \ F$	input of next character (b) to storage location 1	
10	$A \ 1 \ F$	adds $b \times 2^{-16}$ to $C(Acc)$	
11	$S \ 39 \ F$	subtracts 11×2^{-16} from $C(Acc)$	
12	$G \ 4 \ F$	control is transferred to order 4 if $b < 11$	
13	$L \quad D$	multiplies $C(Acc)$ by two $C(Acc) = (2b - 33) 2^{-16}$	an order $A(24+b)F$ is formed if $b > 16$. An order $E(16+b)F$ is formed if $b < 17$. This formed order is placed in 20
14	$S \ 39 \ F$	subtracts 11×2^{-16} from $C(Acc)$ $C(Acc) = (2b - 33) 2^{-16}$	
15	$E \ 17 \ F$	control is transferred to order 17 if the $C(Acc) \geq 0$	
16	$S \ 7 \ F$	subtracts $C(7)$ from $C(Acc)$	
17	$A \ 35 \ F$	adds $C(35)$ to $C(Acc)$ $C(Acc) = E(16+b)F$ if $2b < 33$ $= A(24+b)F$ if $2b > 33$	adds address of order to $C(Acc)$
18	$T \ 20 \ F$	transfers $C(Acc)$ to 20	
19	$A \quad F$	adds the address constructed in storage location 0 to $C(Acc)$	adds selected number to $C(Acc)$ or switches control to order $16+b$
20	$H \ 8 \ F$	puts $C(8)$ equal numerically to 10/32, in multiplier register. This order only used during start it later becomes	
or	$A(24+b)F$ $E(16+b)F$	adds $C(24+b)$ to $C(Acc)$ transfers control to $(16+b)$ as $C(Acc) \geq 0$	adds function digits of order to $C(Acc)$ or during the start adds PD from 40
21	$A \ 40 \ F$	adds $C(40)$ to $C(Acc)$	
22	$T \ 43 \ F$	transfers the order in the accumulator to its place in the store. This order later becomes	assembled order to its location in the store
or	$T \ n \ F$ $E \ m \ F$	switches control to begin obeying the programme—after $EmKPF$	
23	$A \ 22 \ F$	adds $C(22)$ to $C(Acc)$	increases the address specified in order 22 by unity
24	$A \ 2 \ F$	adds $C(2)$ to $C(Acc)$	
25	$T \ 22 \ F$	transfers $C(Acc)$ to 22	control to 34
26	$E \ 34 \ F$	since $C(Acc) = 0$ control is transferred to order 34	
27	$A \ 43 \ F$	adds 2^{-16} to $C(Acc)$	control is switched to order 27 by the order in 20 when the code letter π has been read from the input tape. (π corresponds to 11.) These add 2^{-16} to the address which is in the accumulator and return control to order 8. This order puts back the address in 0, which has thus been increased by 2^{-16}
28	$E \ 8 \ F$	since $C(Acc)$ is positive, control is transferred to order 8	

location of order	order	action of order	general sequence of events
29	A 42 F	adds C (42) to C (Acc)	control switched here by the code letter Z. This adds C (42) to the address in accumulator
30	A 40 F	adds C (40) to C (Acc)	control is switched here by code letter K. The function digits from 40 are now added to C (Acc)
31	E 25 F	if C (Acc) is positive, switch control to order 25; if negative proceed with next order	the accumulator is negative if the code letter K was preceded by G. It is positive if K was preceded by a T or an E, and control is transferred to 25 when the order held in the accumulator is placed in location 22
32	A 22 F	adds C (22) to C (Acc)	the accumulator holds GF at this point and order 22 (TnF) is added to it; as G is the complement of T, PnF is transferred to 42. Note P=0 and so just the address is put in 42
33	T 42 F	transfers C (Acc) to 42, clearing the accumulator	
34	I 40 D	reads next row of holes from the input tape and places them in long location 40	these orders place the function digits in their correct position in 40: control is switched to 8 (see initial orders I). C (35) also used as a constant location 40 now contains the function digits, and 41 contains 0
35	A 40 D	adds C (40') to C (Acc)	
36	R 16 F	shifts C (Acc) six places to the right	
37	T 40 D	transfers C (Acc) to 40' and clears accumulator	
38	E 8 F	since C (Acc) is zero, control is transferred to order 8	constants. PD is transferred to 43 during the start and location 40 thereafter used as working space
39	P 5 D	11×2^{-16}	
40	P D	2^{-16}	

The author wishes to thank Dr M. V. Wilkes and Professor D. R. Hartree for their many helpful suggestions, and also the Department of Scientific and Industrial Research for a maintenance grant.

REFERENCES

- Wilkes, M. V. 1949 *J. Sci. Instrum.* **26**, 217.
 Wilkes, M. V. & Renwick, W. 1949 *J. Sci. Instrum.* **26**, 385.
 Wilkes, M. V. & Renwick, W. 1950 *M.T.A.C.* **4**, 61.

INDEX TO VOLUME 202 (A)

- Adhesion of solids and the effect of surface films (McFarlane & Tabor), 224.
- Aerodynamic drag of grassland (Pasquill), 143.
- Aspects of hybridization (Moffitt), 548.
- Asymptotic expansions for the hypergeometric functions occurring in gas-flow theory (Cherry), 507.
- Atmosphere, lower, turbulence and evaporation (Davies), 96.
- Basu, Sadhan; Sen, Jyotirindra Nath & Palit, Santi R. Degree of polymerization and chain transfer in methyl methacrylate, 485.
- Belcher, H. & Sugden, T. M. Studies on the ionization produced by metallic salts in flames. II. Reactions governed by ionic equilibria in coal-gas/air flames containing alkali metal salts, 17.
- Bhabha, H. J. On the stochastic theory of continuous parametric systems and its application to electron cascades, 301.
- Bloch integral equation and electrical conductivity (Rhodes), 466.
- $\text{CH}_2:\text{CH}.\text{CH}_2-\text{CH}_3$ bond dissociation energy and the heat of formation of the allyl radical (Sehon & Szwarc), 263.
- Chambers, R. G. The conductivity of thin wires in a magnetic field, 378.
- Chang, Su Cheng, Homotopy invariants and continuous mappings, 253.
- Cherry, T. M. Asymptotic expansions for the hypergeometric functions occurring in gas-flow theory, 507.
- Conductivity of thin wires in a magnetic field (Chambers), 378.
- Constitution and mechanism of the selenium rectifier photocell (Preston), 449.
- Contributions to the theory of heat transfer through a laminar boundary layer (Lighthill), 359.
- Copson, E. T. Diffraction by a plane screen, 277.
- Craig, D. P. Electronic levels in simple conjugated systems. I. Configuration interaction in cyclobutadiene, 498.
- Darbyshire, J. Identification of microseismic activity with sea waves, 439.
- Davies, D. R. A note on three-dimensional turbulence and evaporation in the lower atmosphere, 96.
- Degree of polymerization and chain transfer in methyl methacrylate (Basu, Sen and Palit), 485.
- Diffraction by a plane screen (Copson), 277.
- Discovery Committee (Lecture) (Mackintosh), 1.
- EDSAC, programme organization and initial orders (Wheeler), 573.
- Electronic levels in simple conjugated systems. I. Configuration interaction in cyclobutadiene (Craig), 498.
- Fine structure of the hydrogen α line (Kuhn & Series), 127.
- Finite elastic deformation of incompressible isotropic bodies (Green & Shield), 407.
- Finite strains in an anisotropic elastic continuum (Oldroyd), 345.
- Friction and adhesion, relation between (McFarlane & Tabor), 244.
- Grassland, aerodynamic drag of (Pasquill), 143.
- Green, A. E. & Shield, R. T. Finite elastic deformation of incompressible isotropic bodies, 407.
- Hall, G. G. The molecular orbital theory of chemical valency. VI. Properties of equivalent orbitals, 336.
- Hall, G. G. & Lennard-Jones, Sir John. The molecular orbital theory of chemical valency III. Properties of molecular orbitals, 155.
- Heat transfer through a laminar boundary layer (Lighthill), 359.
- Homotopy invariants and continuous mappings (Chang), 253.

- Hybridization (Moffitt), 548.
- Hybrid states (Moffitt), 534.
- Hydrogen α line, fine structure (Kuhn & Series), 127.
- Hydrogen, atomic, interaction with olefines (Melville & Robb), 181.
- Hypergeometric functions in gas-flow theory (Cherry), 507.
- Identification of microseismic activity with sea waves (Darbyshire), 439.
- Instability of liquid surfaces when accelerated in a direction perpendicular to their planes II (Lewis), 81.
- Ionization by metallic salts (Belcher & Sugden), 17.
- Ionosphere, travelling disturbances (Munro), 208.
- Kinetics of the interaction of atomic hydrogen with olefines. V. Results obtained for a further series of compounds (Melville & Robb), 181.
- Kuhn, H. & Series, G. W. The fine structure of the hydrogen α line, 127.
- Le Couteur, K. J. The structure of linear relativistic wave equations. I, 284; II, 394.
- Lennard-Jones, Sir John & Pople, J. A. The molecular orbital theory of chemical valency IV. The significance of equivalent orbitals, 155.
- Lennard-Jones, Sir John. *See* Hall & Lennard-Jones.
- Lewis, D. J. The instability of liquid surfaces when accelerated in a direction perpendicular to their planes II, 81.
- Lighthill, M. J. Contributions to the theory of heat transfer through a laminar boundary layer, 359.
- Liquid surfaces, instability of, when accelerated (Lewis), 81.
- Low temperature resistivity of pure metals (MacDonald & Mendelssohn), 103.
- MacDonald, D. K. C. & Mendelssohn, K. Resistivity of pure metals at low temperatures. I. The alkali metals, 103.
- MacDonald, D. K. C. & Mendelssohn, K. Resistivity of pure metals at low temperatures. II. The alkaline earth metals, 523.
- Mackintosh, N. A. The work of the Discovery Committee, 1.
- McFarlane, J. S. & Tabor, D. Adhesion of solids and the effect of surface films, 224.
- McFarlane, J. S. & Tabor, D. Relation between friction and adhesion, 244.
- Melville, H. W. & Robb, J. C. The kinetics of the interaction of atomic hydrogen with olefines. V. Results obtained for a further series of compounds, 181.
- Mendelssohn, K. *See* MacDonald & Mendelssohn.
- Moffitt, W. Aspects of hybridization, 548.
- Moffitt, W. Term values in hybrid states, 534.
- Molecular orbital theory of chemical valency. III. Properties of molecular orbitals, (Hall & Lennard-Jones), 155.
- Molecular orbital theory of chemical valency. IV. The significance of equivalent orbitals (Lennard-Jones & Pople), 166.
- Molecular orbital theory of chemical valency. V. The structure of water and similar molecules (Pople), 323.
- Molecular orbital theory of chemical valency. VI. Properties of equivalent orbitals (Hall), 336.
- Munro, G. H. Travelling disturbances in the ionosphere, 208.
- Note on three-dimensional turbulence and evaporation in the lower atmosphere (Davies), 96.
- Oldroyd, J. G. Finite strains in an anisotropic elastic continuum, 345.
- Palit, Santi R. *See* Basu, Sen and Palit.
- Pasquill, F. The aerodynamic drag of grassland, 143.
- Polymerization and chain transfer in methyl methacrylate (Basu, Sen & Palit), 485.
- Pople, J. A. *See* Lennard-Jones & Pople.

- Pople, J. A. The molecular orbital theory of chemical valency. V. The structure of water and similar molecules, 323.
- Preston, J. S. Constitution and mechanism of the selenium rectifier photocell, 449.
- Programme organization and initial orders for the EDSAC (Wheeler), 573.
- Quaternary system aluminium-iron-cobalt-nickel, with reference to the role of transitional metals in alloys (Raynor & Waldron), 420.
- Raynor, G. V. & Waldron, M. B. The quaternary system aluminium-iron-cobalt-nickel, with reference to the role of transitional metals in alloys, 420.
- Rhodes, P. The Bloch integral equation and electrical conductivity, 466.
- Resistivity of pure metals at low temperatures. I. The alkali metals (MacDonald & Mendelssohn), 103.
- Resistivity of pure metals at low temperatures. II. The alkaline earth metals (MacDonald & Mendelssohn), 523.
- Robb, J. C. *See* Melville & Robb.
- Sea waves, microseismic activity (Darbyshire), 439.
- Sehon, A. H. & Szwarc, M. $\text{CH}_2:\text{CH}.\text{CH}_3-\text{CH}_2$ bond dissociation energy and the heat of formation of the allyl radical, 263.
- Selenium rectifier photocell (Preston), 449.
- Sen, Jyotirindra Nath. *See* Basu, Nath Sen and Palit.
- Series, G. W. *See* Kuhn & Series.
- Shield, R. T. *See* Green & Shield.
- Statistical mechanics of the two-dimensional assembly (Temperley), 202.
- Stochastic theory of continuous parametric systems and its application to electron cascades (Bhabha), 301.
- Strang, W. J. Transient lift of three-dimensional purely supersonic wings, 54.
- Strang, W. J. Transient source, doublet and vortex solutions of the linearized equations of supersonic flow, 40.
- Structure of linear relativistic wave equations. I. (Le Couteur), 284.
- Structure of linear relativistic waves equations. II. Representations (Le Couteur), 394.
- Studies on the ionization produced by metallic salts in flames. II. Reactions governed by ionic equilibria in coal-gas/air flames containing alkali metal salts (Belcher & Sugden), 17.
- Sugden, T. M. *See* Belcher & Sugden.
- Supersonic flow, linearized equations (Strang), 40.
- Surf beats: sea waves of 1 to 5 min. period (Tucker), 565.
- Szwarc, M. *See* Sehon & Szwarc.
- Tabor, D. *See* McFarlane & Tabor.
- Temperley, H. N. V. Statistical mechanics of the two-dimensional assembly, 202.
- Term values in hybrid states (Moffitt), 534.
- Transient lift of three-dimensional purely supersonic wings (Strang), 54.
- Transient source, doublet and vortex solutions of the linearized equations of supersonic flow (Strang), 40.
- Travelling disturbances in the ionosphere (Munro), 208.
- Tucker, M. J. Surf beats: sea waves of 1 to 5 min. period, 565.
- Valency, chemical, molecular orbital theory (Hall & Lennard-Jones), 155; (Lennard-Jones & Pople), 166; (Pople), 323; (Hall), 336.
- Waldron, M. B. *See* Raynor & Waldron.
- Wave Equations, structure of linear relativistic (Le Couteur), 284.
- Wheeler, D. J. Programme organization and initial orders for the EDSAC, 573.
- Work of the Discovery Committee (Mackintosh), 1.

25 JAN. 1955

